

DreamPRM: Domain-Rewighted Process Reward Model for Multimodal Reasoning

Qi Cao Ruiyi Wang Ruiyi Zhang Sai Ashish Somayajula Pengtao Xie
University of California, San Diego
{q9cao,ruiyi,rz048,ssomayaj,p1xie}@ucsd.edu

Project Page: <https://github.com/coder-qicao/DreamPRM>

Abstract

Reasoning has substantially improved the performance of large language models (LLMs) on complicated tasks. Central to the current reasoning studies, Process Reward Models (PRMs) offer a fine-grained evaluation of intermediate reasoning steps and guide the reasoning process. However, extending PRMs to multimodal large language models (MLLMs) introduces challenges. Since multimodal reasoning covers a wider range of tasks compared to text-only scenarios, the resulting distribution shift from the training to testing sets is more severe, leading to greater generalization difficulty. Training a reliable multimodal PRM, therefore, demands large and diverse datasets to ensure sufficient coverage. However, current multimodal reasoning datasets suffer from a marked quality imbalance, which degrades PRM performance and highlights the need for an effective data selection strategy. To address the issues, we introduce DreamPRM, a domain-reweighted training framework for multimodal PRMs which employs bi-level optimization. In the lower-level optimization, DreamPRM performs fine-tuning on multiple datasets with domain weights, allowing the PRM to prioritize high-quality reasoning signals and alleviating the impact of dataset quality imbalance. In the upper-level optimization, the PRM is evaluated on a separate meta-learning dataset; this feedback updates the domain weights through an aggregation loss function, thereby improving the generalization capability of trained PRM. Extensive experiments on multiple multimodal reasoning benchmarks covering both mathematical and general reasoning show that test-time scaling with DreamPRM consistently improves the performance of state-of-the-art MLLMs. Further comparisons reveal that DreamPRM’s domain-reweighting strategy surpasses other data selection methods and yields higher accuracy gains than existing test-time scaling approaches.

1 Introduction

Reasoning [49] has significantly enhanced the logical and critical thinking capabilities of large language models (LLMs) [1, 7, 52, 43]. Post-training [39, 9] and test-time scaling strategies [38] enable sophisticated reasoning behaviors in LLMs and extend the length of Chain-of-Thoughts (CoTs) [62], thereby achieving strong results on challenging benchmarks [71, 41]. A key component of these advances is the Process Reward Models (PRMs) [26, 24], which provide fine-grained, step-wise supervision of the reasoning process and reliable selection of high-quality reasoning trajectories. These developments are proven highly effective for improving the performance of LLMs in complex tasks [35, 53].

Given the success with LLMs, a natural extension is to apply PRMs to multimodal large language models (MLLMs) [63, 25] to enhance their reasoning abilities. Early studies of multimodal PRMs demonstrate promise results, yet substantial challenges persist. Distinct from text-only inputs of

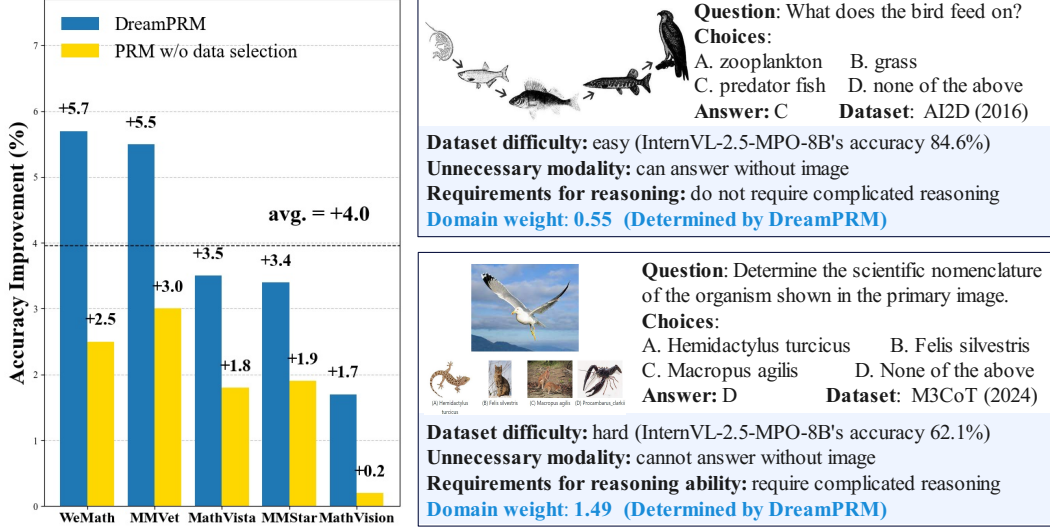


Figure 1: **DreamPRM improves multimodal reasoning by mitigating the dataset quality imbalance problem.** **Left:** On five benchmarks, DreamPRM outperforms base model (InternVL-2.5-8B-MPO [58]) by an average of +4.0%. DreamPRM also consistently surpasses Vanilla PRM trained without data selection. **Right:** Easy AI2D [21] questions (weight 0.55) vs. hard M3CoT [5] questions (weight 1.49) shows how DreamPRM prioritizes data that demand deeper reasoning - samples requiring knowledge from both textual and visual modalities for step-by-step logical deduction.

LLMs, MLLMs must combine diverse visual and language signals: a high-dimensional, continuous image space coupled with discrete language tokens. This fusion dramatically broadens the input manifold and leads to more severe *distribution shifts* [50] from training to testing distributions. Consequently, directly utilizing PRM training strategies from the text domain [60, 34] underperforms, mainly due to the decreased generalizability [10] caused by the insufficient coverage of the multimodal input space.

A straightforward solution to this problem is to combine multiple datasets that emphasize different multimodal reasoning skills, thereby enlarging the sampling space. However, *quality imbalance* among existing multimodal reasoning datasets is more severe than in text-only settings: many contain noisy inputs such as unnecessary modalities [69] or questions of negligible difficulty [30], as illustrated in Fig. 1. Since these easy datasets contribute little to effective sampling, paying much attention to them can substantially degrade PRM performance. Therefore, an effective data selection strategy that filters out unreliable datasets and instances is crucial to training a high-quality multimodal PRM.

To overcome these challenges, we propose DreamPRM, a domain-reweighted training framework for multimodal PRMs. Inspired by domain-reweighting techniques [47, 11, 51], DreamPRM dynamically learns appropriate weights for each multimodal reasoning dataset, allowing them to contribute unequally during training. Datasets that contain many noisy samples tend to receive lower domain weights, reducing their influence on PRM parameter updates. Conversely, high-quality datasets are assigned higher weights and thus play a more important role in optimization. This domain-reweighting strategy alleviates the issue of dataset quality imbalances. DreamPRM adopts a bi-level optimization (BLO) framework [13, 28] to jointly learn the domain weights and PRM parameters. At the lower level, the PRM parameters are optimized with Monte Carlo signals on multiple training domains under different domain weights. At the upper level, the optimized PRM is evaluated on a separate meta domain to compute a novel aggregation function loss, which is used to optimize the domain weights. Extensive experiments on a wide range of multimodal reasoning benchmarks verify the effectiveness of DreamPRM.

Our contributions are summarized as follows:

- We propose DreamPRM, a *domain-reweighted* multimodal process reward model training framework that dynamically adjusts the importance of different training domains. We

formulate the training process of DreamPRM as a *bi-level optimization* (BLO) problem, where the lower level optimizes the PRM via domain-reweighted fine-tuning, and the upper level optimizes domain weights with an aggregation function loss. Our method helps address dataset quality imbalance issue in multimodal reasoning, and improves the generalization ability of PRM.

- We conduct extensive experiments using DreamPRM on a wide range of multimodal reasoning benchmarks. Results indicate that DreamPRM consistently surpasses PRM baselines with other data selection strategies, confirming the effectiveness of its bi-level optimization based domain-reweighting strategy. Carefully designed evaluations further demonstrate that DreamPRM possesses both scaling capability and generalization ability to stronger models.

2 Related Works

Multimodal Reasoning Recent studies have demonstrated that incorporating Chain-of-Thought (CoT) reasoning [61, 22, 72] into LLMs encourages a step-by-step approach, thereby significantly enhancing question-answering performance. However, it has been reported that CoT prompting can’t be easily extended to MLLMs, mainly due to hallucinated outputs during the reasoning process [58, 73, 17]. Therefore, some post-training methods have been proposed for enhancing reasoning capability of MLLMs. InternVL-MPO [58] proposes a mixed preference optimization that jointly optimizes preference ranking, response quality, and response generation loss to improve the reasoning abilities. Llava-CoT [65] creates a structured thinking fine-tuning dataset to make MLLM to perform systematic step-by-step reasoning. Some efforts have also been made for inference time scaling. RLAI-FV [68] proposes a novel self-feedback guidance for inference-time scaling and devises a simple length-normalization strategy tackling the bias towards shorter responses. AR-MCTS [10] combines Monte-Carlo Tree Search (MCTS) and Retrieval Augmented Generation (RAG) to guide MLLM search step by step and explore the answer space.

Process Reward Model Process Reward Model (PRM) [26, 24, 35, 53] provides a more finer-grained verification than Outcome Reward Model (ORM) [8, 46], scoring each step of the reasoning trajectory. However, a central challenge in designing PRMs is obtaining process supervision signals, which require supervised labels for each reasoning step. Current approaches typically depend on costly, labor-intensive human annotation [26], highlighting the need for automated methods to improve scalability and efficiency. Math-Shepherd [55] proposes a method utilizing Monte-Carlo estimation to provide hard labels and soft labels for automatic process supervision. OmegaPRM [34] proposes a Monte Carlo Tree Search (MCTS) for finer-grained exploration for automatic labeling. MiPS [60] further explores the Monte Carlo estimation method and studies the aggregation of PRM signals.

Domain Reweighting Domain reweighting methodologies are developed to modulate the influence of individual data domains, thereby enabling models to achieve robust generalization. Recently, domain reweighting has emerged as a key component in large language model pre-training, where corpora are drawn from heterogeneous sources. DoReMi [64] trains a lightweight proxy model with group distributionally robust optimization to assign domain weights that maximize excess loss relative to a reference model. DOGE [12] proposes a first-order bi-level optimization framework, using gradient alignment between source and target domains to update mixture weights online during training. Complementary to these optimization-based approaches, Data Mixing Laws [67] derives scaling laws that could predict performance under different domain mixtures, enabling low-cost searches for near-optimal weights without proxy models. In this paper, we extend these ideas to process supervision and introduce a novel bi-level domain-reweighting framework.

3 Problem Setting and Preliminaries

Notations. Let $\mathcal{I}$, $\mathcal{T}$, and $\mathcal{Y}$ denote the multimodal input space (images), textual instruction space, and response space, respectively. A multimodal large language model (MLLM) is formalized as a parametric mapping $M_\theta : \mathcal{T} \times \mathcal{I} \rightarrow \Delta(\mathcal{Y})$, where $\hat{y} \sim M_\theta(\cdot|x)$ represents the stochastic generation of responses conditioned on input pair $x = (t, I)$ including visual input $I \in \mathcal{I}$ and textual instruction $t \in \mathcal{T}$, with $\Delta(\mathcal{Y})$ denoting the probability simplex over the response space. We use $y \in \mathcal{Y}$ to denote the ground truth label from a dataset.

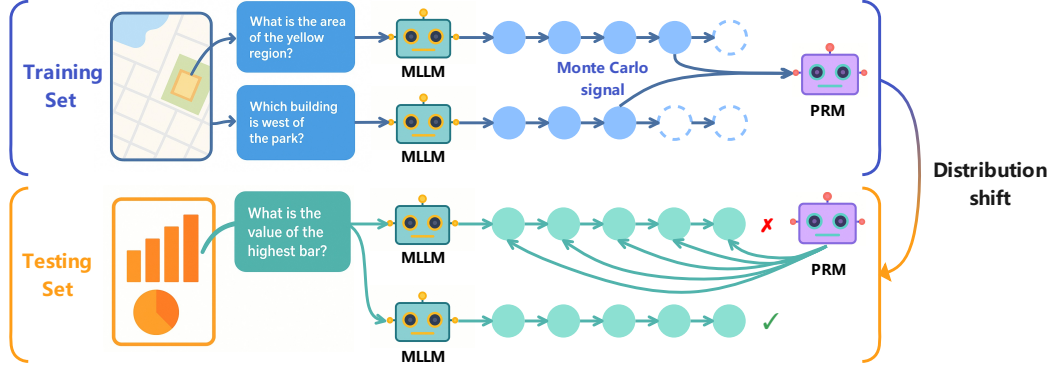


Figure 2: **General flow of training PRM and using PRM for inference.** **Training phase:** Train PRM with Monte Carlo signals from intermediate steps of Chain-of-Thoughts (CoTs). **Inference phase:** Use the trained PRM to verify CoTs step by step and select the best CoT. Conventional training of PRM has poor generalization capability due to *distribution shift* between training set and testing set.

The process reward model (PRM) constitutes a sequence classification function $\mathcal{V}_\phi : \mathcal{T} \times \mathcal{I} \times \mathcal{Y} \rightarrow [0, 1]$, parameterized by ϕ , which quantifies the epistemic value of partial reasoning state $\hat{y}_i$ through scalar reward $p_i = \mathcal{V}_\phi(x, \hat{y}_i)$, modeling incremental utility toward solving instruction t under visual grounding I . Specifically, $\hat{y}_i$ represents the first i steps of a complete reasoning trajectory $\hat{y}$.

PRM training with Monte Carlo signals. Due to the lack of ground truth epistemic value for each partial reasoning state $\hat{y}_i$, training of PRM requires automatic generation of approximated supervision signals. An effective approach to obtain these signals is to use the Monte Carlo method [60, 56]. We first feed the input question-image pair $x = (t, I)$ and the prefix solution $\hat{y}_i$ into the MLLM, and let it complete the remaining steps until reaching the final answer. We randomly sample multiple completions, compare their final answers to the gold answer y , and thereby obtain multiple correctness labels. PRM is trained as a sequence classification task to predict these correctness labels. The ratio of correct completions at the i -th step estimates the “correctness level” up to step i , which is used as the approximated supervision signals p_i to train the PRM. Formally,

$$p_i = \text{MonteCarlo}(x, \hat{y}_i, y) = \frac{\text{num}(\text{correct completions from } \hat{y}_i)}{\text{num}(\text{total completions from } \hat{y}_i)} \quad (1)$$

PRM-based inference with aggregation function. After training a PRM, a typical way of conducting PRM-based MLLM inference is to use aggregation function [60]. Specifically, for each candidate solution $\hat{y}$ from the MLLM, PRM will generate a list of predicted probabilities $p = \{p_1, p_2, \dots, p_n\}$ accordingly, one for each step $\hat{y}_i$ in the solution. The list of predicted probabilities are then aggregated using the following function:

$$\mathcal{A}(p) = \sum_{i=1}^n \log \frac{p_i}{1 - p_i}. \quad (2)$$

The aggregated value corresponds to the score of a specific prediction $\hat{y}$, and the final PRM-based solution is the one with the highest aggregated score.

Bi-level optimization. Bi-level optimization (BLO) has been widely used in meta-learning [13], neural architecture search [28], and data reweighting [48]. A BLO problem is usually formulated as:

$$\min_{\alpha} \mathcal{U}(\alpha, \phi^*(\alpha)) \quad (3)$$

$$s.t. \phi^*(\alpha) = \arg \min_{\phi} \mathcal{L}(\phi, \alpha) \quad (4)$$

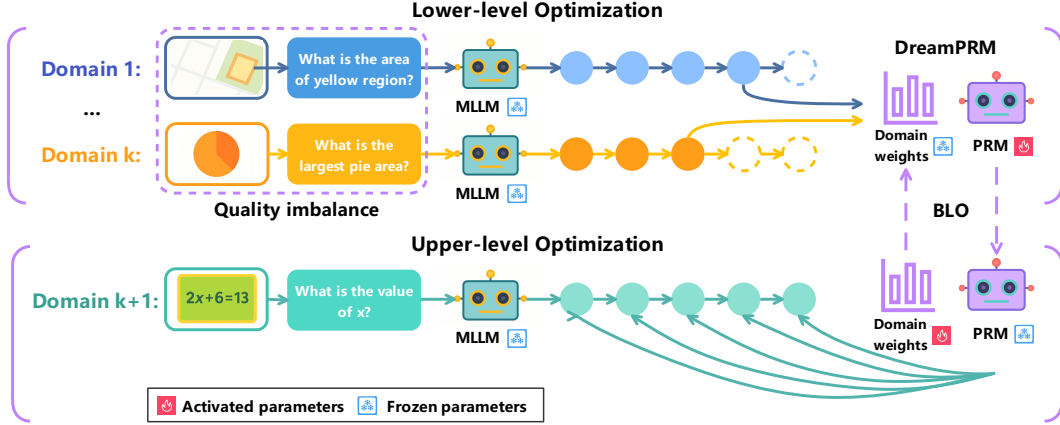


Figure 3: **The proposed bi-level optimization based domain-reweighting method.** **Lower-level optimization:** In this stage, PRM’s parameters are updated on multiple datasets with domain weights, allowing the PRM to prioritize domains with better quality. **Upper-level optimization:** In this stage, the PRM is evaluated on a separate meta dataset to compute an aggregation function loss and optimize the domain weights. DreamPRM helps address dataset *quality imbalance* problems and leads to stronger and more generalizable reasoning performance.

where $\mathcal{U}$ is the upper-level optimization problem (OP) with parameter α , and $\mathcal{L}$ is the lower-level OP with parameter ϕ . The lower-level OP is nested within the upper-level one, and the two OPs are mutually dependent.

4 The Proposed Domain-reweighting Method

Overview. Training process reward models (PRMs) for MLLMs is challenging for two reasons: (1) dataset (domain) quality imbalance, and (2) discrepancy between training and inference procedures. To address these two challenges, we propose DreamPRM, which automatically searches for domain importance using a novel aggregation function loss that better simulates the inference process of PRM. Under a bi-level optimization framework, it optimizes PRM parameters with Monte Carlo signals at the lower level, and optimizes trainable domain importance weights with aggregation function loss at the upper level. An overview of DreamPRM method is shown in Fig. 3.

Datasets. We begin with $K+1$ datasets, each from a distinct domain (e.g., science, geometry). The first K datasets form the training pool $\mathcal{D}_{tr} = \{\mathcal{D}_1, \dots, \mathcal{D}_K\}$, while the remaining dataset, $\mathcal{D}_{meta} = \mathcal{D}_{K+1}$, is a meta (validation) dataset with better quality.

Lower-level optimization: domain-reweighted training of PRM. In lower-level optimization, we aim to update the weights ϕ of PRM with domain-reweighted training. We first define the typical PRM training loss $\mathcal{L}_{tr}$ on a single domain $\mathcal{D}_k$, given PRM parameters ϕ , as follows:

$$\mathcal{L}_{tr}(\mathcal{D}_k, \phi) = \sum_{(x,y) \in \mathcal{D}_k} \sum_{i=1}^n \mathcal{L}_{MSE}(\mathcal{V}_\phi(x, \hat{y}_i), p_i) \quad (5)$$

where $\hat{y}_i$ is the prefix of MLLM generated text $\hat{y} = M_\theta(x)$ given input pair $x = (t, I)$, and p_i is the process supervision signal value obtained by Monte Carlo estimation given input pair x , prefix $\hat{y}_i$ and ground truth label y , as previously defined in Equation 1. The PRM is optimized by minimizing the mean squared error (MSE) between supervision signal and PRM predicted score $\mathcal{V}_\phi(x, \hat{y}_i)$. With the PRM training loss on a single domain $\mathcal{D}_k$ above, we next define the domain-reweighted training objective of PRM on multiple training domains $\mathcal{D} = \{\mathcal{D}_k\}_{k=1}^K$. The overall objective is a weighted sum of the single-domain PRM training losses, allowing the contribution of each domain to be adjusted during the learning process:

$$\mathcal{L}_{tr}(\mathcal{D}_{tr}, \phi, \alpha) = \sum_{k=1}^K \alpha_k \mathcal{L}_{tr}(\mathcal{D}_k, \phi) \quad (6)$$

Here, $\alpha = \{\alpha_k\}_{k=1}^K$ represents the trainable domain weight parameters, indicating the importance of each domain. By optimizing this objective, we obtain the optimal value of PRM parameters ϕ^* :

$$\phi^*(\alpha) = \arg \min_{\phi} \mathcal{L}_{tr}(\mathcal{D}_{tr}, \phi, \alpha) \quad (7)$$

It is worth mentioning that only ϕ is optimized at this level, while α remains fixed.

Upper-level optimization: learning domain reweighting parameters. In upper-level optimization, we optimize the domain reweighting parameter α on meta dataset $\mathcal{D}_{meta}$ given optimal PRM weights $\phi^*(\alpha)$ obtained from the lower level. To make the meta learning target more closely reflect the actual PRM-based inference process, we propose a novel meta loss function $\mathcal{L}_{meta}$, different from the training loss $\mathcal{L}_{tr}$. Specifically, we first obtain an aggregated score $\mathcal{A}(p)$ for each generated solution $\hat{y}$ from the MLLM given input pair $x = (t, I)$, following process in Section 3. We then create a ground truth signal $r(\hat{y}, y)$ by assigning it a value of 1 if the generated $\hat{y}$ contains ground truth y , and 0 otherwise. The meta loss is defined as the mean squared error between aggregated score and ground truth signal:

$$\mathcal{L}_{meta}(\mathcal{D}_{meta}, \phi^*(\alpha)) = \sum_{(x, y) \in \mathcal{D}_{meta}} \mathcal{L}_{MSE}(\sigma(\mathcal{A}(\mathcal{V}_{\phi^*(\alpha)}(x, \hat{y}))), r(\hat{y}, y)) \quad (8)$$

where $\mathcal{A}$ represents the aggregation function as previously defined in Equation 2, and σ denotes the sigmoid function to map the aggregated score to a probability. Accordingly, the optimization problem at the upper level is formulated as follows:

$$\min_{\alpha} \mathcal{L}_{meta}(\mathcal{D}_{meta}, \phi^*(\alpha)) \quad (9)$$

To solve this optimization problem, we propose an efficient gradient-based algorithm, which is detailed in Appendix A.

5 Experimental Results

5.1 Datasets and benchmarks

For datasets used in lower-level optimization ($\mathcal{D}_{tr}$ in Section 4), our study utilizes a diverse set of datasets spanning multiple domains to ensure comprehensive coverage of multimodal reasoning tasks. The selected 15 multimodal datasets cover 4 major categories, including science, chart, geometry and commonsense, with a wide range of task types (e.g., visual question answering, optical character recognition, spatial understanding). Details of these datasets are provided in Appendix B.

For the dataset used in upper-level optimization ($\mathcal{D}_{meta}$ in Section 4), we select data from MMMU [70] to simulate a realistic and diverse reasoning scenario. MMMU focuses on advanced perception and reasoning with domain-specific knowledge. Its questions span 30 subjects and 183 subfields, comprising 30 highly heterogeneous image types, such as charts, diagrams, maps, tables, music sheets, and chemical structures.

At evaluation time, we use five multimodal reasoning benchmarks for testing the capability of DreamPRM. WEMATH [42], MATHVISTA [30], and MATHVISION [54] focus more on math-related reasoning tasks and logic and critical thinking, while MMVET [69] and MMSTAR [4] focus more on real-life tasks that require common knowledge and general reasoning abilities.

Table 1: **Comparative evaluation of DreamPRM and baselines on multimodal reasoning benchmarks.** **Bold numbers** indicate the best performance, while underlined numbers indicate the second best. The table reports accuracy (%) on five datasets: WEMATH, MATHVISTA, MATHVISION, MMVET, and MMSTAR.

	Math Reasoning			General Reasoning	
	WEMATH (loose)	MATHVISTA (testmini)	MATHVISION (test)	MMVET (v1)	MMSTAR (test)
<i>Zero-shot Methods</i>					
Gemini-1.5-Pro [44]	46.0	63.9	19.2	<u>64.0</u>	59.1
GPT-4v [40]	51.4	49.9	<u>21.7</u>	67.7	<u>62.0</u>
LLaVA-OneVision-7B [23]	44.8	63.2	18.4	57.5	61.7
Qwen2-VL-7B [57]	42.9	58.2	16.3	62.0	60.7
InternVL-2.5-8B-MPO [58]	51.7	65.4	20.4	55.9	58.9
<i>Test-time Scaling Methods (InternVL-2.5-8B-MPO based)</i>					
Self-consistency [59]	56.4	67.1	20.7	57.4	59.6
Self-correction [16]	54.0	63.8	21.6	54.9	59.7
ORM [46]	56.9	65.3	20.5	55.9	60.1
Vanilla PRM [26]	54.2	67.2	20.6	58.9	60.8
CaR-PRM [15]	54.7	<u>67.5</u>	21.0	60.6	61.1
s1-PRM [38]	<u>57.1</u>	65.8	20.2	60.1	60.4
DreamPRM (ours)	57.4	68.9	22.1	61.4	62.3

5.2 Experimental settings

Multistage reasoning. To elicit consistent steady reasoning responses from current MLLMs, we draw on the Llava-CoT approach [66], which fosters structured thinking prior to answer generation. Specifically, we prompt MLLMs to follow five reasoning steps: (1) Restate the question. (2) Gather evidence from the image. (3) Identify any background knowledge needed. (4) Reason with the current evidence. (5) Summarize and conclude with all the information. We also explore zero-shot prompting settings in conjunction with structural reasoning, which can be found in Appendix C. We use 8 different chain-of-thought reasoning trajectories for all test-time scaling methods, unless otherwise stated.

Base models. For inference, we use InternVL-2.5-8B-MPO [58] as the base MLLM, which has undergone post-training to enhance its reasoning abilities and is well-suited for our experiment. For fine-tuning PRM, we adopt Qwen2-VL-2B-Instruct [57]. Qwen2-VL is a state-of-the-art multimodal model pretrained for general vision-language understanding tasks. This pretrained model serves as the initialization for our fine-tuning process.

Training hyperparameters. In the lower-level optimization, we perform 5 inner gradient steps per outer update (unroll steps = 5) using the AdamW [29] optimizer with learning rate set to 5×10^{-7} . In the upper-level optimization, we use the AdamW optimizer (lr = 0.01, weight decay = 10^{-3}) and a StepLR scheduler (step size = 5000, $\gamma = 0.5$). In total, DreamPRM is fine-tuned for 10000 iterations. Our method is implemented with Betty [6], and the fine-tuning process takes approximately 10 hours on two NVIDIA A100 GPUs.

Baselines. We use three major categories of baselines: (1) State-of-the-art models on public leaderboards, including Gemini-1.5-Pro [44], GPT-4V [40], LLaVA-OneVision-7B [23], Qwen2-VL-7B [57]. We also carefully reproduce the results of InternVL-2.5-8B-MPO with structural thinking. (2) Test-time scaling methods (excluding PRM) based on the InternVL-2.5-8B-MPO model, including: (i) Self-consistency [59], which selects the most consistent reasoning chain via majority voting over multiple responses; (ii) Self-correction [16], which prompts the model to critically reflect

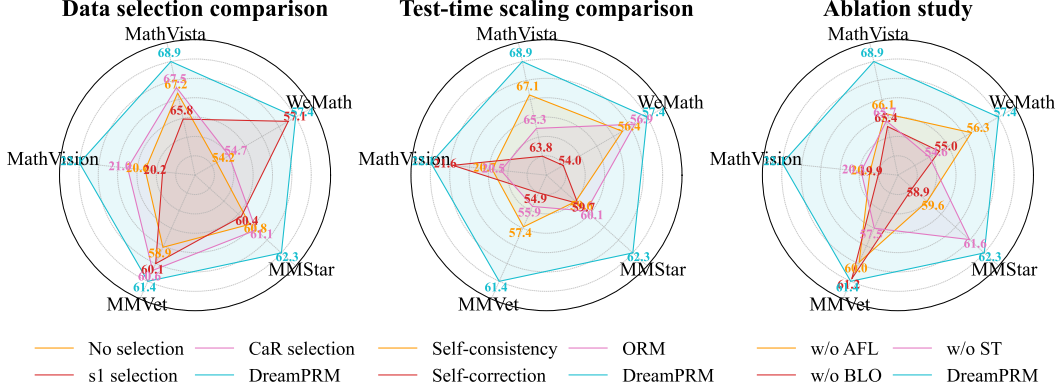


Figure 4: **Comparative evaluation of DreamPRM on multimodal reasoning benchmarks.** Radar charts report accuracy (%) on five datasets (WEMATH, MATHVISTA, MATHVISION, MMVET, and MMSTAR). (a) Impact of different data selection strategies. (b) Comparison with existing test-time scaling methods. (c) Ablation study of three key components, i.e. w/o aggregation function loss (AFL), w/o bi-level optimization (BLO), and w/o structural thinking (ST).

on and revise its initial answers; and (iii) Outcome Reward Model (ORM) [46], which evaluates and scores the final response to select the most promising one. (3) PRM-based methods, including: (i) Vanilla PRM trained without any data selection, as commonly used in LLM settings [26]; (ii) s1-PRM, which selects high-quality reasoning responses based on three criteria - difficulty, quality, and diversity - following the s1 strategy [38]; and (iii) CaR-PRM, which filters high-quality visual questions using clustering and ranking techniques, as proposed in CaR [15].

5.3 Benchmark evaluation of DreamPRM

Tab. 1 presents the primary experimental results. We observe that: (1) **DreamPRM outperforms other PRM-based methods**, highlighting the effectiveness of our domain reweighting strategy. Compared to the vanilla PRM trained without any data selection, DreamPRM achieves a consistent performance gain of 2%-3% across all five datasets, suggesting that effective data selection is crucial for training high-quality multimodal PRMs. Moreover, DreamPRM also outperforms s1-PRM and CaR-PRM, which rely on manually designed heuristic rules for data selection. These results indicate that selecting suitable reasoning datasets for PRM training is a complex task, and handcrafted rules are often suboptimal. In contrast, our automatic domain-reweighting approach enables the model to adaptively optimize its learning process, illustrating how data-driven optimization offers a scalable solution to dataset selection challenges. (2) **DreamPRM outperforms SOTA MLLMs with much fewer parameters**, highlighting the effectiveness of DreamPRM. For example, DreamPRM significantly surpasses two trillion-scale closed-source LLMs (GPT-4v and Gemini-1.5-Pro) on 4 out of 5 datasets. In addition, it consistently improves the performance of the base model, InternVL-2.5-8B-MPO, achieving an average gain of 4% on the five datasets. These results confirm that DreamPRM effectively yields a high-quality PRM, which is capable of enhancing multimodal reasoning across a wide range of benchmarks. (3) **DreamPRM outperforms other test-time scaling methods**, primarily because it enables the training of a high-quality PRM that conducts fine-grained, step-level evaluation. While most test-time scaling methods yield moderate improvements, DreamPRM leads to the most substantial gains, suggesting that the quality of the reward model is critical for effective test-time scaling. We further provide case studies in Appendix D, which intuitively illustrate how DreamPRM assigns higher scores to coherent and high-quality reasoning trajectories.

5.4 Scaling and generalization analysis of DreamPRM

DreamPRM scales reliably with more CoT candidates. As shown in the left panel of Fig. 5, the accuracy of DreamPRM consistently improves on all five benchmarks as the number of CoTs increases from $k=2$ to $k=8$, expanding the radar plot outward. Intuitively, a larger set of candidates increases the likelihood of including high-quality reasoning trajectories, but it also makes identifying the best ones more challenging. The consistent performance gains indicate that DreamPRM effectively

Table 2: Accuracy on MathVista using DreamPRM with varying numbers k of CoTs.

Model Name	pass@1	DreamPRM (select k CoTs)			
	$k=1$	$k=2$	$k=4$	$k=6$	$k=8$
InternVL-2.5-8B-MPO [58]	65.4	65.3	66.5	67.8	68.9
GPT-4.1-mini (4-14-25) [40]	71.5	71.8	72.5	73.2	74.4
o4-mini (4-16-25) [39]	80.6	81.6	82.5	84.2	85.2

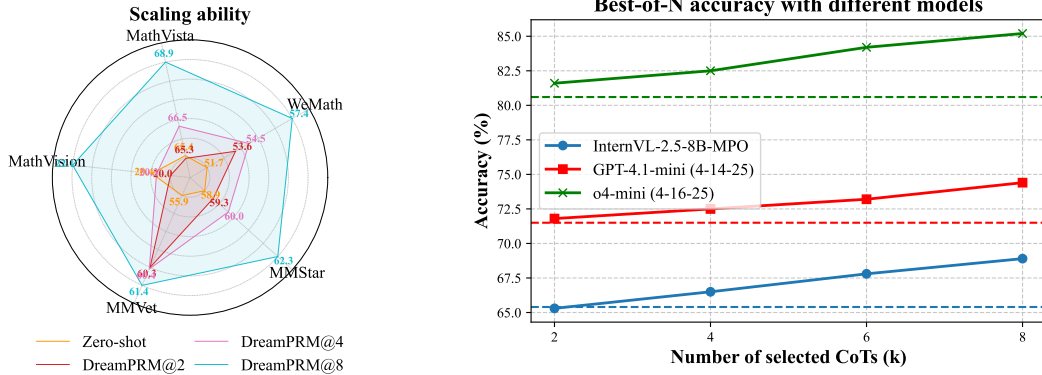


Figure 5: **Scaling ability and cross-model generalization.** (a) Radar chart of five multimodal reasoning benchmarks shows that DreamPRM delivers monotonic accuracy gains as the number of selected chains-of-thought increases (@2, @4, @8) over the pass@1 baseline. (b) Best-of- N accuracy curves for InternVL-2.5-8B-MPO (blue) and GPT-4.1-mini (red) on MATHVISTA confirm that the same DreamPRM-ranked CoTs generalize across models, consistently outperforming pass@1 performance (dashed lines) as k grows.

verifies and ranks CoTs, demonstrating its robustness in selecting high-quality reasoning trajectories under more complex candidate pools.

DreamPRM transfers seamlessly to stronger base MLLMs. The right panel of Fig.5 shows the MATHVISTA accuracy when applying DreamPRM to recent MLLMs, GPT-4.1-mini (2025-04-14) [40] and o4-mini (2025-04-16) [39]. For o4-mini model, the pass@1 score of 80.6% steadily increases to 85.2% at $k=8$, surpassing the previous state-of-the-art performance. This best-of- N trend, previously observed with InternVL, also holds for GPT-4.1-mini and o4-mini, demonstrating the generalization ability of DreamPRM. Full results of these experiments are provided in Tab. 2.

5.5 Ablation study

In this section, we investigate the importance of three components in DreamPRM: (1) bi-level optimization, (2) aggregation function loss in upper-level, and (3) structural thinking prompt (detailed in Section 5.2). As shown in the rightmost panel of Fig. 4, the complete DreamPRM achieves the best results compared to three ablation baselines across all five benchmarks. Eliminating bi-level optimization causes large performance drop (e.g., -3.5% on MATHVISTA and -3.4% on MMSTAR). Removing aggregation function loss leads to a consistent 1%-2% decline (e.g., 57.4% $\rightarrow$ 56.3% on

Table 3: Ablation study evaluating the impact of individual components of DreamPRM

Ablation / Dataset	WEMATH	MATHVISTA	MATHVISION	MMVET	MMSTAR
DreamPRM (original)	57.4	68.9	22.1	61.4	62.3
w/o aggregation function loss	56.3 (-1.1)	66.1 (-2.8)	20.1 (-2.0)	60.0 (-1.4)	59.6 (-2.7)
w/o bi-level optimization	55.0 (-2.4)	65.4 (-3.5)	19.9 (-2.2)	61.2 (-0.2)	58.9 (-3.4)
w/o structural thinking	54.6 (-2.8)	65.7 (-3.2)	20.3 (-1.8)	57.5 (-3.9)	61.6 (-0.7)

WEMATH). Excluding structural thinking also degrades performance (e.g., -1.8% on MATHVISION). These results indicate that all three components are critical for DreamPRM to achieve the best performance. More detailed results are shown in Tab. 3.

5.6 Analysis of learned domain weights

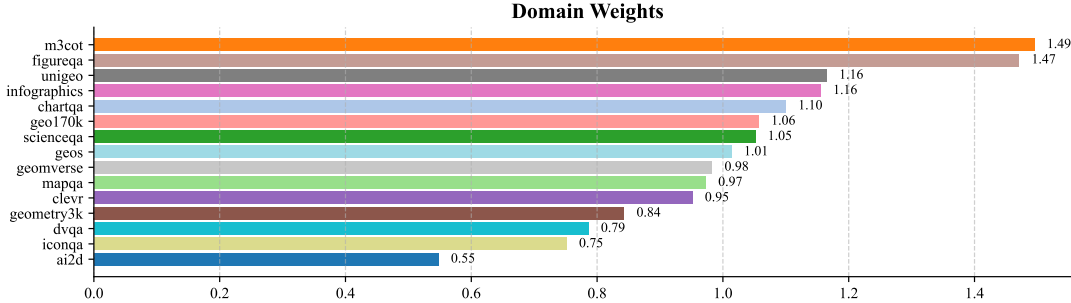


Figure 6: Learned domain weights after the convergence of the DreamPRM training process.

The final domain weights (Fig. 6) range from 0.55 to 1.49: M3CoT [5] and FIGUREQA [19] receive the highest weights (approximately 1.5), while AI2D [21] and ICONQA [33] are assigned lower weights (less than 0.8). This learned weighting pattern contributes to improved PRM performance, indicating that the quality imbalance problem across reasoning datasets is real and consequential. Additionally, as shown in Fig. 8 in Appendix, all domain weights are initialized to 1.0 and eventually converge during the training process of DreamPRM.

6 Conclusions

We propose DreamPRM, the first domain-reweighted PRM framework for multimodal reasoning. By automatically searching for domain weights using a bi-level optimization framework, DreamPRM effectively mitigates issues caused by dataset quality imbalance and significantly enhances the generalizability of multimodal PRMs. Extensive experiments on five diverse benchmarks confirm that DreamPRM outperforms both vanilla PRMs without domain reweighting and PRMs using heuristic data selection methods. We also observe that the domain weights learned by DreamPRM correlate with dataset quality, effectively separating challenging, informative sources from overly simplistic or noisy ones. These results highlight the effectiveness of our proposed automatic domain reweighting strategy.

References

- [1] Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners, 2020.
- [2] Shuaichen Chang, David Palzer, Jialin Li, Eric Fosler-Lussier, and Ningchuan Xiao. Mapqa: A dataset for question answering on choropleth maps, 2022.
- [3] Jiaqi Chen, Tong Li, Jinghui Qin, Pan Lu, Liang Lin, Chongyu Chen, and Xiaodan Liang. Unigeo: Unifying geometry logical reasoning via reformulating mathematical expression, 2022.
- [4] Lin Chen, Jinsong Li, Xiaoyi Dong, Pan Zhang, Yuhang Zang, Zehui Chen, Haodong Duan, Jiaqi Wang, Yu Qiao, Dahua Lin, and Feng Zhao. Are we on the right way for evaluating large vision-language models?, 2024.
- [5] Qiguang Chen, Libo Qin, Jin Zhang, Zhi Chen, Xiao Xu, and Wanxiang Che. M³cot: A novel benchmark for multi-domain multi-step multi-modal chain-of-thought, 2024.

- [6] Sang Keun Choe, Willie Neiswanger, Pengtao Xie, and Eric Xing. Betty: An automatic differentiation library for multilevel optimization. In *The Eleventh International Conference on Learning Representations*, 2023.
- [7] Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, Parker Schuh, Kensen Shi, Sasha Tsvyashchenko, Joshua Maynez, Abhishek Rao, Parker Barnes, Yi Tay, Noam Shazeer, Vinodkumar Prabhakaran, Emily Reif, Nan Du, Ben Hutchinson, Reiner Pope, James Bradbury, Jacob Austin, Michael Isard, Guy Gur-Ari, Pengcheng Yin, Toju Duke, Anselm Levskaya, Sanjay Ghemawat, Sunipa Dev, Henryk Michalewski, Xavier Garcia, Vedant Misra, Kevin Robinson, Liam Fedus, Denny Zhou, Daphne Ippolito, David Luan, Hyeontaek Lim, Barret Zoph, Alexander Spiridonov, Ryan Sepassi, David Dohan, Shrivani Agrawal, Mark Omernick, Andrew M. Dai, Thanumalayan Sankaranarayanan Pillai, Marie Pellat, Aitor Lewkowycz, Erica Moreira, Rewon Child, Oleksandr Polozov, Katherine Lee, Zongwei Zhou, Xuezhi Wang, Brennan Saeta, Mark Diaz, Orhan Firat, Michele Catasta, Jason Wei, Kathy Meier-Hellstern, Douglas Eck, Jeff Dean, Slav Petrov, and Noah Fiedel. Palm: Scaling language modeling with pathways, 2022.
- [8] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training verifiers to solve math word problems, 2021.
- [9] DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jia Shi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojuan Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanjia Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhen Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu Zhang, and Zhen Zhang. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, 2025.
- [10] Guanting Dong, Chenghao Zhang, Mengjie Deng, Yutao Zhu, Zhicheng Dou, and Ji-Rong Wen. Progressive multimodal reasoning via active retrieval, 2024.
- [11] Simin Fan, Matteo Pagliardini, and Martin Jaggi. Doge: Domain reweighting with generalization estimation, 2024.
- [12] Simin Fan, Matteo Pagliardini, and Martin Jaggi. DOGE: Domain reweighting with generalization estimation. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 12895–12915. PMLR, 21–27 Jul 2024.
- [13] Chelsea Finn, P. Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *International Conference on Machine Learning*, 2017.
- [14] Jiahui Gao, Renjie Pi, Jipeng Zhang, Jiacheng Ye, Wanjuan Zhong, Yufei Wang, Lanqing Hong, Jianhua Han, Hang Xu, Zhenguo Li, and Lingpeng Kong. G-llava: Solving geometric problem with multi-modal large language model, 2023.

- [15] Yuan Ge, Yilun Liu, Chi Hu, Weibin Meng, Shimin Tao, Xiaofeng Zhao, Hongxia Ma, Li Zhang, Boxing Chen, Hao Yang, Bei Li, Tong Xiao, and Jingbo Zhu. Clustering and ranking: Diversity-preserved instruction selection through expert-aligned quality estimation, 2024.
- [16] Jiayi He, Hehai Lin, Qingyun Wang, Yi Fung, and Heng Ji. Self-correction is more than refinement: A learning framework for visual and language reasoning tasks, 2024.
- [17] Dongzhi Jiang, Renrui Zhang, Ziyu Guo, Yanwei Li, Yu Qi, Xinyan Chen, Liuhui Wang, Jianhan Jin, Claire Guo, Shen Yan, Bo Zhang, Chaoyou Fu, Peng Gao, and Hongsheng Li. Mme-cot: Benchmarking chain-of-thought in large multimodal models for reasoning quality, robustness, and efficiency, 2025.
- [18] Kushal Kafle, Brian Price, Scott Cohen, and Christopher Kanan. Dvqa: Understanding data visualizations via question answering, 2018.
- [19] Samira Ebrahimi Kahou, Vincent Michalski, Adam Atkinson, Akos Kadar, Adam Trischler, and Yoshua Bengio. Figureqa: An annotated figure dataset for visual reasoning, 2018.
- [20] Mehran Kazemi, Hamidreza Alvari, Ankit Anand, Jialin Wu, Xi Chen, and Radu Soricut. Geomverse: A systematic evaluation of large models for geometric reasoning, 2023.
- [21] Aniruddha Kembhavi, Mike Salvato, Eric Kolve, Minjoon Seo, Hannaneh Hajishirzi, and Ali Farhadi. A diagram is worth a dozen images, 2016.
- [22] Takeshi Kojima, Shixiang (Shane) Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. Large language models are zero-shot reasoners. In *Advances in Neural Information Processing Systems*, volume 35, pages 22199–22213, 2022.
- [23] Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng Li, Hao Zhang, Kaichen Zhang, Peiyuan Zhang, Yanwei Li, Ziwei Liu, and Chunyuan Li. Llava-onevision: Easy visual task transfer, 2024.
- [24] Yifei Li, Zeqi Lin, Shizhuo Zhang, Qiang Fu, Bei Chen, Jian-Guang Lou, and Weizhu Chen. Making large language models better reasoners with step-aware verifier, 2023.
- [25] Zongxia Li, Xiyang Wu, Hongyang Du, Fuxiao Liu, Huy Nghiem, and Guangyao Shi. A survey of state of the art large vision language models: Alignment, benchmark, evaluations and challenges. *arXiv preprint arXiv:2501.02189*, 2025.
- [26] Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. In *The Twelfth International Conference on Learning Representations*, 2024.
- [27] Adam Dahlgren Lindström and Savitha Sam Abraham. Clevr-math: A dataset for compositional language, visual and mathematical reasoning, 2022.
- [28] Hanxiao Liu, Karen Simonyan, and Yiming Yang. DARTS: Differentiable architecture search. In *International Conference on Learning Representations*, 2019.
- [29] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization, 2019.
- [30] Pan Lu, Hritik Bansal, Tony Xia, Jiacheng Liu, Chunyuan Li, Hannaneh Hajishirzi, Hao Cheng, Kai-Wei Chang, Michel Galley, and Jianfeng Gao. Mathvista: Evaluating mathematical reasoning of foundation models in visual contexts. In *International Conference on Learning Representations (ICLR)*, 2024.
- [31] Pan Lu, Ran Gong, Shibiao Jiang, Liang Qiu, Siyuan Huang, Xiaodan Liang, and Song-Chun Zhu. Inter-gps: Interpretable geometry problem solving with formal language and symbolic reasoning. In *The 59th Annual Meeting of the Association for Computational Linguistics (ACL)*, 2021.
- [32] Pan Lu, Swaroop Mishra, Tony Xia, Liang Qiu, Kai-Wei Chang, Song-Chun Zhu, Oyvind Tafjord, Peter Clark, and Ashwin Kalyan. Learn to explain: Multimodal reasoning via thought chains for science question answering, 2022.
- [33] Pan Lu, Liang Qiu, Jiaqi Chen, Tony Xia, Yizhou Zhao, Wei Zhang, Zhou Yu, Xiaodan Liang, and Song-Chun Zhu. Iconqa: A new benchmark for abstract diagram understanding and visual language reasoning, 2022.
- [34] Liangchen Luo, Yinxiao Liu, Rosanne Liu, Samrat Phatale, Meiqi Guo, Harsh Lara, Yunxuan Li, Lei Shu, Yun Zhu, Lei Meng, Jiao Sun, and Abhinav Rastogi. Improve mathematical reasoning in language models by automated process supervision, 2024.

- [35] Qianli Ma, Haotian Zhou, Tingkai Liu, Jianbo Yuan, Pengfei Liu, Yang You, and Hongxia Yang. Let’s reward step by step: Step-level reward model as the navigators for reasoning, 2023.
- [36] Ahmed Masry, Do Xuan Long, Jia Qing Tan, Shafiq Joty, and Enamul Hoque. Chartqa: A benchmark for question answering about charts with visual and logical reasoning, 2022.
- [37] Minesh Mathew, Viraj Bagal, Rubèn Pérez Tito, Dimosthenis Karatzas, Ernest Valveny, and C. V Jawahar. Infographicvqa, 2021.
- [38] Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. s1: Simple test-time scaling, 2025.
- [39] OpenAI, :, Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, Alex Iftimie, Alex Karpenko, Alex Tachard Passos, Alexander Neitz, Alexander Prokofiev, Alexander Wei, Allison Tam, Ally Bennett, Ananya Kumar, Andre Saraiva, Andrea Vallone, Andrew Duberstein, Andrew Kondrich, Andrey Mishchenko, Andy Applebaum, Angela Jiang, Ashvin Nair, Barret Zoph, Behrooz Ghorbani, Ben Rossen, Benjamin Sokolowsky, Boaz Barak, Bob McGrew, Borys Minaiev, Botao Hao, Bowen Baker, Brandon Houghton, Brandon McKinzie, Brydon Eastman, Camillo Lugaresi, Cary Bassin, Cary Hudson, Chak Ming Li, Charles de Bourcy, Chelsea Voss, Chen Shen, Chong Zhang, Chris Koch, Chris Orsinger, Christopher Hesse, Claudia Fischer, Clive Chan, Dan Roberts, Daniel Kappler, Daniel Levy, Daniel Selsam, David Dohan, David Farhi, David Mely, David Robinson, Dimitris Tsipras, Doug Li, Dragos Oprica, Eben Freeman, Eddie Zhang, Edmund Wong, Elizabeth Proehl, Enoch Cheung, Eric Mitchell, Eric Wallace, Erik Ritter, Evan Mays, Fan Wang, Felipe Petroski Such, Filippo Raso, Florencia Leoni, Foivos Tsimpourlas, Francis Song, Fred von Lohmann, Freddie Sulit, Geoff Salmon, Giambattista Parascandolo, Gildas Chabot, Grace Zhao, Greg Brockman, Guillaume Leclerc, Hadi Salman, Haiming Bao, Hao Sheng, Hart Andrin, Hessam Bagherinezhad, Hongyu Ren, Hunter Lightman, Hyung Won Chung, Ian Kivlichan, Ian O’Connell, Ian Osband, Ignasi Clavera Gilaberte, Ilge Akkaya, Ilya Kostrikov, Ilya Sutskever, Irina Kofman, Jakub Pachocki, James Lennon, Jason Wei, Jean Harb, Jerry Twore, Jiacheng Feng, Jiahui Yu, Jiayi Weng, Jie Tang, Jieqi Yu, Joaquin Quiñero Candela, Joe Palermo, Joel Parish, Johannes Heidecke, John Hallman, John Rizzo, Jonathan Gordon, Jonathan Uesato, Jonathan Ward, Joost Huizinga, Julie Wang, Kai Chen, Kai Xiao, Karan Singhal, Karina Nguyen, Karl Cobbe, Katy Shi, Kayla Wood, Kendra Rimbach, Keren Gu-Lemberg, Kevin Liu, Kevin Lu, Kevin Stone, Kevin Yu, Lama Ahmad, Lauren Yang, Leo Liu, Leon Maksin, Leyton Ho, Liam Fedus, Lilian Weng, Linden Li, Lindsay McCallum, Lindsey Held, Lorenz Kuhn, Lukas Kondraciuk, Lukasz Kaiser, Luke Metz, Madelaine Boyd, Maja Trebacz, Manas Joglekar, Mark Chen, Marko Tintor, Mason Meyer, Matt Jones, Matt Kaufer, Max Schwarzer, Meghan Shah, Mehmet Yatbaz, Melody Y. Guan, Mengyuan Xu, Mengyuan Yan, Mia Glaese, Mianna Chen, Michael Lampe, Michael Malek, Michele Wang, Michelle Fradin, Mike McClay, Mikhail Pavlov, Miles Wang, Mingxuan Wang, Mira Murati, Mo Bavarian, Mostafa Rohaninejad, Nat McAleese, Neil Chowdhury, Neil Chowdhury, Nick Ryder, Nikolas Tezak, Noam Brown, Ofir Nachum, Oleg Boiko, Oleg Murk, Olivia Watkins, Patrick Chao, Paul Ashbourne, Pavel Izmailov, Peter Zhokhov, Rachel Dias, Rahul Arora, Randall Lin, Rapha Gontijo Lopes, Raz Gaon, Reah Miyara, Reimar Leike, Renny Hwang, Rhythm Garg, Robin Brown, Roshan James, Rui Shu, Ryan Cheu, Ryan Greene, Saachi Jain, Sam Altman, Sam Toizer, Sam Toyer, Samuel Miserendino, Sandhini Agarwal, Santiago Hernandez, Sasha Baker, Scott McKinney, Scottie Yan, Shengjia Zhao, Shengli Hu, Shibani Santurkar, Shraman Ray Chaudhuri, Shuyuan Zhang, Siyuan Fu, Spencer Papay, Steph Lin, Suchir Balaji, Suvansh Sanjeev, Szymon Sidor, Tal Broda, Aidan Clark, Tao Wang, Taylor Gordon, Ted Sanders, Tejal Patwardhan, Thibault Sottiaux, Thomas Degry, Thomas Dimson, Tianhao Zheng, Timur Garipov, Tom Stasi, Trapit Bansal, Trevor Creech, Troy Peterson, Tyna Eloundou, Valerie Qi, Vineet Kosaraju, Vinnie Monaco, Vitchyr Pong, Vlad Fomenko, Weiye Zheng, Wenda Zhou, Wes McCabe, Wojciech Zaremba, Yann Dubois, Yinghai Lu, Yining Chen, Young Cha, Yu Bai, Yuchen He, Yuchen Zhang, Yunyun Wang, Zheng Shao, and Zhuohan Li. Openai o1 system card, 2024.
- [40] OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, Red Avila, Igor Babuschkin, Suchir Balaji, Valerie Balcom, Paul Baltescu, Haiming Bao, Mohammad Bavarian, Jeff Belgum, Irwan Bello, Jake Berdine, Gabriel Bernadett-Shapiro, Christopher Berner, Lenny Bogdonoff, Oleg Boiko, Madelaine Boyd, Anna-Luisa Brakman, Greg Brockman, Tim Brooks, Miles Brundage, Kevin Button, Trevor Cai, Rosie Campbell, Andrew Cann, Brittany Carey, Chelsea Carlson, Rory Carmichael, Brooke Chan, Che Chang, Fotis Chantzis, Derek Chen, Sully Chen, Ruby Chen, Jason Chen, Mark Chen, Ben Chess, Chester Cho, Casey Chu, Hyung Won Chung, Dave Cummings, Jeremiah Currier, Yunxing Dai, Cory Decareaux, Thomas Degry, Noah Deutsch, Damien Deville, Arka Dhar, David Dohan, Steve Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti, Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix, Simón Posada Fishman, Juston Forte, Isabella Fulford, Leo Gao, Elie Georges, Christian Gibson, Vik Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-Lopes, Jonathan Gordon, Morgan Grafstein, Scott Gray, Ryan Greene, Joshua Gross, Shixiang Shane Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris, Yuchen

He, Mike Heaton, Johannes Heidecke, Chris Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele, Brandon Houghton, Kenny Hsu, Shengli Hu, Xin Hu, Joost Huizinga, Shantanu Jain, Shawn Jain, Joanne Jang, Angela Jiang, Roger Jiang, Haozhun Jin, Denny Jin, Shino Jomoto, Billie Jonn, Heewoo Jun, Tomer Kaftan, Łukasz Kaiser, Ali Kamali, Ingmar Kanitscheider, Nitish Shirish Keskar, Tabarak Khan, Logan Kilpatrick, Jong Wook Kim, Christina Kim, Yongjik Kim, Jan Hendrik Kirchner, Jamie Kiros, Matt Knight, Daniel Kokotajlo, Łukasz Kondraciuk, Andrew Kondrich, Aris Konstantinidis, Kyle Kosic, Gretchen Krueger, Vishal Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan Leike, Jade Leung, Daniel Levy, Chak Ming Li, Rachel Lim, Molly Lin, Stephanie Lin, Mateusz Litwin, Theresa Lopez, Ryan Lowe, Patricia Lue, Anna Makanju, Kim Malfacini, Sam Manning, Todor Markov, Yaniv Markovski, Bianca Martin, Katie Mayer, Andrew Mayne, Bob McGrew, Scott Mayer McKinney, Christine McLeavey, Paul McMillan, Jake McNeil, David Medina, Aalok Mehta, Jacob Menick, Luke Metz, Andrey Mishchenko, Pamela Mishkin, Vinnie Monaco, Evan Morikawa, Daniel Mossing, Tong Mu, Mira Murati, Oleg Murk, David Mély, Ashvin Nair, Reiichiro Nakano, Rajeev Nayak, Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh, Long Ouyang, Cullen O’Keefe, Jakub Pachocki, Alex Paino, Joe Palermo, Ashley Pantuliano, Giambattista Parascandolo, Joel Parish, Emy Parparita, Alex Passos, Mikhail Pavlov, Andrew Peng, Adam Perelman, Filipe de Avila Belbute Peres, Michael Petrov, Henrique Ponde de Oliveira Pinto, Michael, Pokorny, Michelle Pocrass, Vitchyr H. Pong, Tolly Powell, Alethea Power, Boris Power, Elizabeth Proehl, Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh, Cameron Raymond, Francis Real, Kendra Rimbach, Carl Ross, Bob Rotsted, Henri Roussez, Nick Ryder, Mario Saltarelli, Ted Sanders, Shibani Santurkar, Girish Sastry, Heather Schmidt, David Schnurr, John Schulman, Daniel Selsam, Kyla Sheppard, Toki Sherbakov, Jessica Shieh, Sarah Shoker, Pranav Shyam, Szymon Sidor, Eric Sigler, Maddie Simens, Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin Sokolowsky, Yang Song, Natalie Staudacher, Felipe Petroski Such, Natalie Summers, Ilya Sutskever, Jie Tang, Nikolas Tezak, Madeleine B. Thompson, Phil Tillet, Amin Tootoonchian, Elizabeth Tseng, Preston Tuggle, Nick Turley, Jerry Tworek, Juan Felipe Cerón Uribe, Andrea Vallone, Arun Vijayvergiya, Chelsea Voss, Carroll Wainwright, Justin Jay Wang, Alvin Wang, Ben Wang, Jonathan Ward, Jason Wei, CJ Weinmann, Akila Welihinda, Peter Welinder, Jiayi Weng, Lilian Weng, Matt Wiethoff, Dave Willner, Clemens Winter, Samuel Wolrich, Hannah Wong, Lauren Workman, Sherwin Wu, Jeff Wu, Michael Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qiming Yuan, Wojciech Zaremba, Rowan Zellers, Chong Zhang, Marvin Zhang, Shengjia Zhao, Tianhao Zheng, Juntang Zhuang, William Zhuk, and Barret Zoph. Gpt-4 technical report, 2024.

- [41] Guilherme Penedo, Anton Lozhkov, Hynek Kydlíček, Loubna Ben Allal, Edward Beeching, Agustín Piqueres Lajarín, Quentin Gallouédec, Nathan Habib, Lewis Tunstall, and Leandro von Werra. Codeforces. <https://huggingface.co/datasets/open-r1/codeforces>, 2025.
- [42] Runqi Qiao, Qiuna Tan, Guanting Dong, Minhui Wu, Chong Sun, Xiaoshuai Song, Zhuoma GongQue, Shanglin Lei, Zhe Wei, Miaoxuan Zhang, Runfeng Qiao, Yifan Zhang, Xiao Zong, Yida Xu, Muxi Diao, Zhimin Bao, Chen Li, and Honggang Zhang. We-math: Does your large multimodal model achieve human-like mathematical reasoning?, 2024.
- [43] Qwen, :, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report, 2025.
- [44] Alex Reid et al. Gemini 1.5: Unlocking multimodal understanding across millions of tokens, 2024.
- [45] Minjoon Seo, Hannaneh Hajishirzi, Ali Farhadi, Oren Etzioni, and Clint Malcolm. Solving geometry problems: Combining text and diagram interpretation. In Lluís Màrquez, Chris Callison-Burch, and Jian Su, editors, *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, pages 1466–1476, Lisbon, Portugal, September 2015. Association for Computational Linguistics.
- [46] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. Deepseekmath: Pushing the limits of mathematical reasoning in open language models, 2024.
- [47] Jun Shu, Qi Xie, Lixuan Yi, Qian Zhao, Sanping Zhou, Zongben Xu, and Deyu Meng. Meta-weight-net: Learning an explicit mapping for sample weighting, 2019.
- [48] Jun Shu, Qi Xie, Lixuan Yi, Qian Zhao, Sanping Zhou, Zongben Xu, and Deyu Meng. Meta-weight-net: Learning an explicit mapping for sample weighting. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d’Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.

- [49] Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling llm test-time compute optimally can be more effective than scaling model parameters, 2024.
- [50] Shezheng Song, Xiaopeng Li, Shasha Li, Shan Zhao, Jie Yu, Jun Ma, Xiaoguang Mao, and Weimin Zhang. How to bridge the gap between modalities: Survey on multimodal large language model, 2025.
- [51] Daouda Sow, Herbert Woitschläger, Saikiran Bulusu, Shiqiang Wang, Hans-Arno Jacobsen, and Yingbin Liang. Dynamic loss-based sample reweighting for improved large language model pretraining, 2025.
- [52] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. Llama: Open and efficient foundation language models, 2023.
- [53] Chaojie Wang, Yanchen Deng, Zhiyi Lyu, Liang Zeng, Jujie He, Shuicheng Yan, and Bo An. Q*: Improving multi-step reasoning for llms with deliberative planning, 2024.
- [54] Ke Wang, Junting Pan, Weikang Shi, Zimu Lu, Mingjie Zhan, and Hongsheng Li. Measuring multimodal mathematical reasoning with math-vision dataset, 2024.
- [55] Peiyi Wang, Lei Li, Zhihong Shao, Runxin Xu, Damai Dai, Yifei Li, Deli Chen, Yu Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce LLMs step-by-step without human annotations. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar, editors, *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9426–9439, Bangkok, Thailand, August 2024. Association for Computational Linguistics.
- [56] Peiyi Wang, Lei Li, Zhihong Shao, Runxin Xu, Damai Dai, Yifei Li, Deli Chen, Yu Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce LLMs step-by-step without human annotations. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar, editors, *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9426–9439, Bangkok, Thailand, August 2024. Association for Computational Linguistics.
- [57] Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Yang Fan, Kai Dang, Mengfei Du, Xuancheng Ren, Rui Men, Dayiheng Liu, Chang Zhou, Jingren Zhou, and Junyang Lin. Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. *arXiv preprint arXiv:2409.12191*, 2024.
- [58] Weiyun Wang, Zhe Chen, Wenhai Wang, Yue Cao, Yangzhou Liu, Zhangwei Gao, Jinguo Zhu, Xizhou Zhu, Lewei Lu, Yu Qiao, and Jifeng Dai. Enhancing the reasoning ability of multimodal large language models via mixed preference optimization. *arXiv preprint arXiv:2411.10442*, 2024.
- [59] Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. Self-consistency improves chain of thought reasoning in language models, 2023.
- [60] Zihan Wang, Yunxuan Li, Yuexin Wu, Liangchen Luo, Le Hou, Hongkun Yu, and Jingbo Shang. Multi-step problem solving through a verifier: An empirical analysis on model-induced process supervision, 2024.
- [61] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, brian ichter, Fei Xia, Ed H. Chi, Quoc V Le, and Denny Zhou. Chain of thought prompting elicits reasoning in large language models. In Alice H. Oh, Alekh Agarwal, Danielle Belgrave, and Kyunghyun Cho, editors, *Advances in Neural Information Processing Systems*, 2022.
- [62] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models. In *Advances in Neural Information Processing Systems*, volume 35, pages 24824–24837. Curran Associates, Inc., 2022.
- [63] Jiayang Wu, Wensheng Gan, Zefeng Chen, Shicheng Wan, and Philip S. Yu. Multimodal large language models: A survey, 2023.
- [64] Sang Michael Xie, Hieu Pham, Xuanyi Dong, Nan Du, Hanxiao Liu, Yifeng Lu, Percy Liang, Quoc V Le, Tengyu Ma, and Adams Wei Yu. Doremi: Optimizing data mixtures speeds up language model pretraining. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [65] Guowei Xu, Peng Jin, Hao Li, Yibing Song, Lichao Sun, and Li Yuan. Llava-cot: Let vision language models reason step-by-step, 2024.
- [66] Guowei Xu, Peng Jin, Hao Li, Yibing Song, Lichao Sun, and Li Yuan. Llava-cot: Let vision language models reason step-by-step, 2025.

- [67] Jiasheng Ye, Peiju Liu, Tianxiang Sun, Jun Zhan, Yunhua Zhou, and Xipeng Qiu. Data mixing laws: Optimizing data mixtures by predicting language modeling performance. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [68] Tianyu Yu, Haoye Zhang, Qiming Li, Qixin Xu, Yuan Yao, Da Chen, Xiaoman Lu, Ganqu Cui, Yunkai Dang, Taiwen He, Xiaocheng Feng, Jun Song, Bo Zheng, Zhiyuan Liu, Tat-Seng Chua, and Maosong Sun. Rlaif-v: Open-source ai feedback leads to super gpt-4v trustworthiness. *arXiv preprint arXiv:2405.17220*, 2024.
- [69] Weihao Yu, Zhengyuan Yang, Linjie Li, Jianfeng Wang, Kevin Lin, Zicheng Liu, Xinchao Wang, and Lijuan Wang. Mm-vet: Evaluating large multimodal models for integrated capabilities, 2024.
- [70] Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruoqi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, Cong Wei, Botao Yu, Ruibin Yuan, Renliang Sun, Ming Yin, Boyuan Zheng, Zhenzhu Yang, Yibo Liu, Wenhao Huang, Huan Sun, Yu Su, and Wenhui Chen. Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi, 2024.
- [71] Di Zhang. Aime_1983_2024 (revision 6283828), 2025.
- [72] Zhuosheng Zhang, Aston Zhang, Mu Li, and Alex Smola. Automatic chain of thought prompting in large language models. In *The Eleventh International Conference on Learning Representations*, 2023.
- [73] Haojie Zheng, Tianyang Xu, Hanchi Sun, Shu Pu, Ruoxi Chen, and Lichao Sun. Thinking before looking: Improving multimodal llm reasoning via mitigating visual hallucination, 2024.

A Optimization algorithm

Directly solving the bi-level optimization problem in Equation 9 can be computational prohibitive due to its nested structure. Following previous work [6], we use approximated algorithm with a few unrolling steps. For example, under one-step unrolling, the updating of PRM’s weights can be expressed as:

$$\phi^{(t+1)} = \phi^{(t)} - \beta_1 \nabla_{\phi} \mathcal{L}_{tr}(\mathcal{D}_{tr}, \phi, \alpha) \quad (10)$$

where β_1 is the learning rate in lower level optimization. After obtaining the updated PRM parameter $\phi^{(t+1)}$ from Equation 10, the domain-reweighting parameter α is then updated as follows:

$$\alpha^{(t+1)} = \alpha^{(t)} - \beta_2 \nabla_{\alpha} \mathcal{L}_{meta}(\mathcal{D}_{meta}, \phi^*(\alpha)) \quad (11)$$

where β_2 is the learning rate for upper level optimization. The two optimization steps in Equation 10 and Equation 11 are conducted iteratively until convergence to get optimal PRM weights ϕ^* and optimal domain reweighting parameter α^* .

B Dataset Details

Table 4: Multimodal datasets involved in the fine-tuning of DreamPRM, organized by task category.

Task	Dataset
Science	AI2D [21], ScienceQA [32], M3CoT [5]
Chart	ChartQA [36], DVQA [18], MapQA [2], FigureQA [19]
Geometry	Geo170k [14], Geometry3K [31], UniGeo [3], GeomVerse [20], GeoS [45]
Commonsense	IconQA [33], InfographicsVQA [37], CLEVR-Math [27]

To enhance the robustness of DreamPRM, we collect a diverse set of datasets in lower-level optimization, spanning multiple domains to ensure a comprehensive coverage of multimodal reasoning tasks, as reported in Tab. 4. The selected 15 multimodal datasets covers 4 major categories including science, chart, geometry and commonsense, with a wide range of task types (QA, OCR, spatial understanding). Additionally, we observe that for some questions, given the current structural thinking prompts, MLLMs consistently produce either correct or incorrect answers. Continuing to sample such questions is a waste of computational resources. Inspired by the dynamic sampling strategy in DAPO [69], we propose a similar dynamic sampling technique for Monte Carlo estimation that focuses on prompts with varied outcomes to improve efficiency. After processing and sampling, the training datasets in lower-level $\mathcal{D}_{tr}$ have around 15k examples (1k per each of the 15 domains), while the meta dataset in the upper-level $\mathcal{D}_{meta}$ has around 1k validation examples from the MMMU [70] dataset.

C Structural Thinking Prompt

The detailed structural thinking prompt applied in our experiments is reported in Fig. 7. We carefully design 5 reasoning steps to boost the reasoning capabilities of the MLLMs and enable process supervision.

D Additional Experimental Results

Loss curves and domain weights. The loss curves and domain weights during the fine-tuning of DreamPRM are illustrated in Fig. 8. It can be observed that the learnt distribution emphasizes informative mathematical figure domains while attenuating less relevant sources. Additionally, domain weights start at 1.0 and quickly diverge, stabilizing after roughly half the training, and the

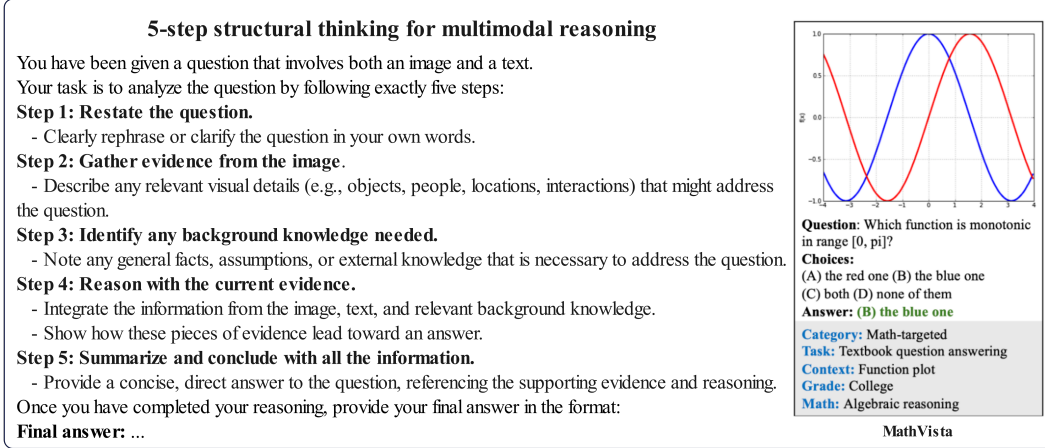


Figure 7: Zero-shot prompting for structural thinking.

inner and outer losses decrease steadily and plateau, indicating stable convergence of the bi-level training procedure.

Case study. A complete case study illustrating DreamPRM’s step-wise evaluation is reported in Fig. 9. DreamPRM assigns higher scores to high-quality, coherent reasoning steps, while penalizes flawed or unsupported steps.

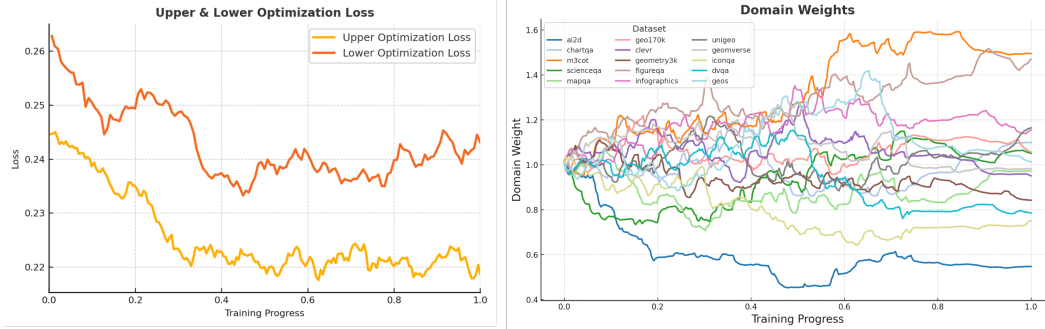
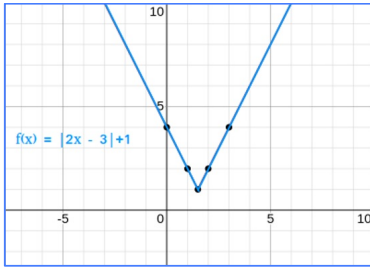


Figure 8: Optimization loss curves and dynamic domain weights throughout DreamPRM fine-tuning.

E Limitations & Future Work.

DreamPRM currently assumes a fixed set of domains and requires Monte-Carlo sampling, which can be computationally heavy. Future work could explore instance-level reweighting, adaptive sampling strategies, and integration with retrieval-augmented generation to further cut compute while broadening coverage. We will release code, trained weights, and evaluation scripts to facilitate reproducibility and community adoption.



Question:

The derivative of $f(x)$ at $x=2$ is ____ that at $x=5$

Choices:

(A) larger than (B) equal to (C) smaller than

Ground Truth: B

Benchmark: MathVista

Step 1: Restate the question.

- The question is asking whether the derivative of the function $f(x) = |2x - 3| + 1$ at $x = 2$ is larger than, equal to, or smaller than the derivative at $x = 5$. [DreamPRM: 0.628]

Step 2: Gather evidence from the image.

- The image shows the graph of the function $f(x) = |2x - 3| + 1$.
- The graph has a V-shape, indicating it is an absolute value function. [DreamPRM: 0.575]

Step 3: Identify any background knowledge needed.

- The derivative of an absolute value function changes at the point where the expression inside the absolute value is zero.
- For $f(x) = |2x - 3| + 1$, the expression inside the absolute value is $2x - 3$.
- The derivative of $|2x - 3| + 1$ is 2 for $x > 1.5$ and -2 for $x < 1.5$. [DreamPRM: 0.598]

Step 4: Reason with the current evidence.

- The point where the derivative changes is at $x = 1.5$.
- For $x = 2$, the function is in the interval where $x > 1.5$, so the derivative is 2 .
- For $x = 5$, the function is also in the interval where $x > 1.5$, so the derivative is 2 . [DreamPRM: 0.748]

Step 5: Summarize and conclude with all the information.

- Both the derivatives at $x = 2$ and $x = 5$ are 2 , so they are equal. [DreamPRM: 0.812]

Final answer: B ✓

Step 1: Restate the question.

- The question is asking how the derivative of $f(x)$ at $x = 2$ compares to the derivative of $f(x)$ at $x = 5$. [DreamPRM: 0.676]

Step 2: Gather evidence from the image.

- The image shows the graph of the function $f(x) = |2x - 3| + 1$.
- There are points marked on the graph at $x = 2$ and $x = 5$. [DreamPRM: 0.499]

Step 3: Identify any background knowledge needed.

- For absolute value functions, the derivative involves identifying where the function switches from decreasing to increasing.
- The function $f(x) = |2x - 3| + 1$ changes slope at the vertex of the absolute value, which is at $x = \frac{3}{2}$. [DreamPRM: 0.561]

Step 4: Reason with the current evidence.

- At $x = 2$, the function appears linear with a negative slope (since it's to the left of the vertex).
- At $x = 5$, the function appears linear with a positive slope (since it's to the right of the vertex).
- The slope at $x = 2$ (negative) is steeper/more negative than the slope at $x = 5$ (positive). [DreamPRM: 0.397]

Step 5: Summarize and conclude with all the information.

- The derivative at $x = 2$ is negative and steeper than the derivative at $x = 5$, which is positive.
- Therefore, the derivative at $x = 2$ is larger in absolute value than the derivative at $x = 5$. [DreamPRM: 0.396]

Final answer: A ✗

Figure 9: A case study of DreamPRM's step-wise evaluation.