

OpenMMReasoner: Pushing the Frontiers for Multimodal Reasoning with an Open and General Recipe

Kaichen Zhang^{*,1,2,4}, Keming Wu^{*,1,3,4}, Zuhao Yang^{1,2,4}, Bo Li^{2,4}, Kairui Hu^{2,4},
Bin Wang³, Ziwei Liu^{2,4}, Xingxuan Li¹, Lidong Bing¹

¹MiroMind AI, ²Nanyang Technological University, ³Tsinghua University, ⁴LMMs-Lab Team

Email Contact: {zhan0564}@e.ntu.edu.sg, {xingxuan.li, lidong.bing}@miromind.ai

<https://evolvinglmm-lab.github.io/OpenMMReasoner/>

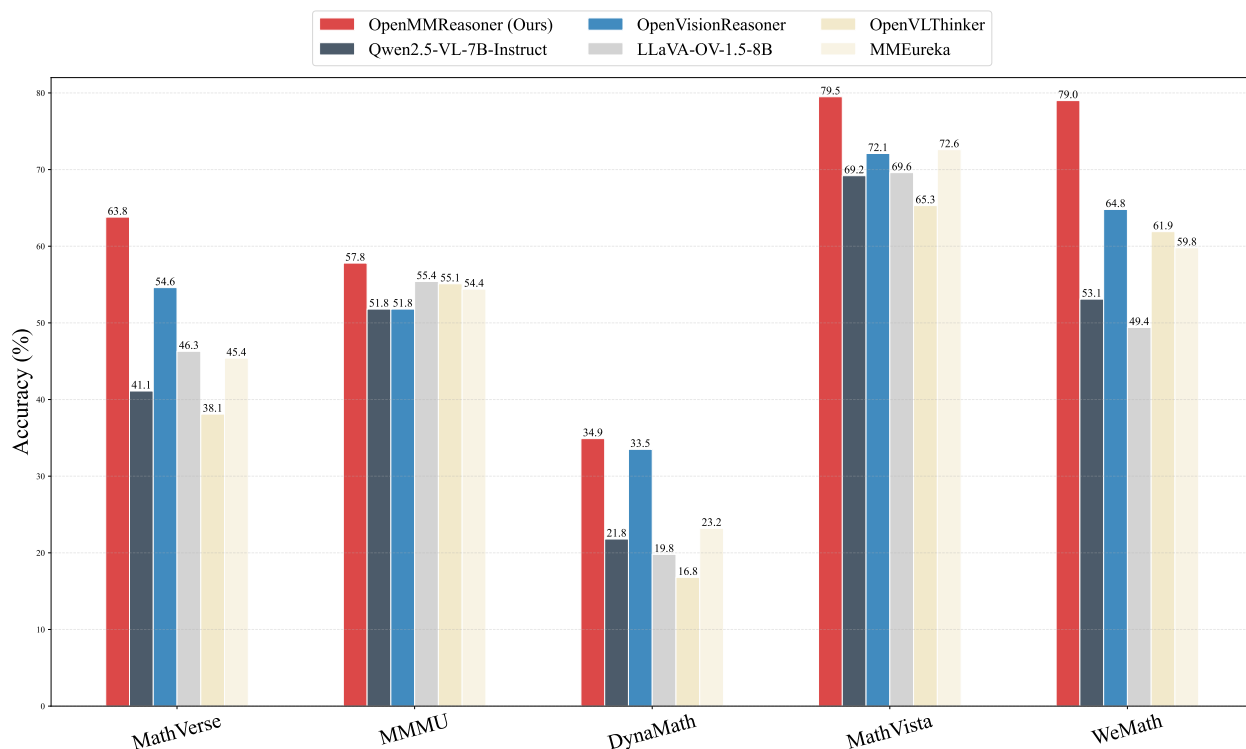


Figure 1. **Performance Comparison with State-of-the-Art Large Multimodal Reasoning Models across Various Benchmarks.** Our proposed **OpenMMReasoner** consistently outperforms competing methods, highlighting its effectiveness in complex reasoning tasks.

Abstract

Recent advancements in large reasoning models have fueled growing interest in extending such capabilities to multimodal domains. However, despite notable progress in visual reasoning, the lack of transparent and reproducible data curation and training strategies remains a major barrier to scalable research. In this work, we introduce **OpenMMReasoner**, a fully transparent two-stage recipe for multimodal reasoning spanning supervised fine-tuning (SFT) and reinforcement learning (RL). In the SFT stage, we construct

an 874K-sample cold-start dataset with rigorous step-by-step validation, providing a strong foundation for reasoning capabilities. The subsequent RL stage leverages a 74K-sample dataset across diverse domains to further sharpen and stabilize these abilities, resulting in a more robust and efficient learning process. Extensive evaluations demonstrate that our training recipe not only surpasses strong baselines but also highlights the critical role of data quality and training design in shaping multimodal reasoning performance. Notably, our method achieves a 11.6% improvement over the *Qwen2.5-VL-7B-Instruct* baseline across nine multimodal reasoning benchmarks, establishing a solid

*Equal contribution.  Corresponding author.

empirical foundation for future large-scale multimodal reasoning research. We open-sourced all our codes, pipeline, and data at <https://github.com/EvolvingLLMs-Lab/OpenMMReasoner>.

1. Introduction

With the rapid progress of reinforcement learning with verifiable rewards (RLVR), leading models such as DeepSeek-R1 [11] and OpenAI o3 [26] have demonstrated strong reasoning capabilities across mathematics, programming, and scientific tasks. These advancements highlight RLVR as an effective paradigm for improving structured reasoning in large language models (LLMs). Motivated by this progress, researchers have increasingly explored RL-enhanced reasoning from LLMs to large multimodal models (LMMs). Recent studies [4, 8, 25, 38] push the boundaries of multimodal reasoning, showing that reinforcement learning (RL) can strengthen fine-grained visual understanding and complex cross-modal problem-solving. These developments suggest that the benefits of RLVR can transfer beyond text, offering a scalable path toward more capable and reliable multimodal reasoning systems.

Despite these advancements, transparency across the training pipeline remains limited. While numerous studies investigate the supervised-finetuning (SFT) and RL stages [3, 7, 26], few provide details of their data curation processes or conduct comprehensive ablation analyses. This lack of openness restricts reproducibility and obscures a deeper understanding of how reasoning-capable LMMs are actually built and how their training dynamics evolve. Recent efforts such as ThinkLite-VL [36] and MM-Eureka [25] propose clearer methodologies for RL data construction and offer useful insights for algorithm design. Other work, including Open Vision Reasoner (OVR) [38], attempts to address both SFT and RL, but does not provide a scalable, unified recipe that generalizes across tasks and modalities. Similarly, LLM-focused efforts [10] emphasize SFT while offering limited discussion of RL or multimodal extensions. Collectively, these limitations point to a critical gap: the lack of a transparent, scalable, and training recipe that unifies SFT and RL for building multimodal reasoning systems.

In this work, we present a comprehensive and scalable empirical study on training Large Multimodal Reasoning Models (LMRMs) from open-source LLMs (e.g., Qwen2.5-VL [3]). Through extensive experimentation, we develop simple yet effective data pipelines for constructing high-quality SFT and RL datasets, enabling reliable cold-start training of a strong LMRM. Our analysis provides detailed insights into which components most effectively contribute at each stage, revealing practical principles for building robust, generalizable, and efficient multimodal reasoning systems. Furthermore, we explore how to maximize scaling efficiency across both SFT and RL phases and summarize

key lessons learned throughout the process.

Building on these insights, we introduce **OpenMMReasoner**, a fully transparent training recipe for multimodal reasoning. It features (1) a high-quality 874k-sample SFT dataset that establishes a strong reasoning foundation prior to RL, and (2) a 74k-sample RL dataset accompanied by detailed analysis of diverse designs and optimization strategies, which further sharpens and stabilizes these capabilities. As shown in Fig. 1, our recipe yields a highly capable LMRM that consistently outperforms state-of-the-art methods such as OVR across a wide range of multimodal reasoning benchmarks, demonstrating both the effectiveness and scalability of our approach. The complete pipeline, spanning both the SFT and RL phases, is illustrated in Fig. 2. As summarized in Tab. 1, all components of our workflow are fully open and transparent, providing a comprehensive and reproducible view of the entire data curation and training process.

To summarize, our contributions are threefold: **(1) Comprehensive Data Curation Insights.** We present the first systematic study on curating high-quality SFT and RL data for multimodal reasoning, supported by extensive and rigorous experiments across diverse modalities and reasoning types. We found that scaling data diversity is a critical factor for curating high-quality datasets. While diversity in data sources is important, diversity in answers represents an additional essential axis for improvement. **(2) A Strong SFT Recipe for Reasoning.** We present a robust and reproducible SFT recipe that effectively equips LMMs with strong reasoning capabilities. Our approach incorporates step-by-step validation and offers practical insights for scaling a high-quality data pipeline focused on reasoning. By carefully selecting an appropriate teacher model for rejection sampling, and incorporating cross-domain data sources, we construct a cold-start dataset that exhibits both high diversity and quality, forming a solid reasoning foundation for subsequent RL. **(3) An Advanced RL Recipe for Reasoning Enhancement.** We conduct a comprehensive comparative analysis of multiple RL strategies, including GSPO [53], GRPO [31], and DAPO [45], to evaluate their stability, efficiency, and scaling behavior. By selecting the most suitable algorithm, we establish a robust RL pipeline that further sharpens and stabilizes reasoning abilities, delivering both high stability and superior performance.

2. Related Work

RL has emerged as a powerful approach for enhancing reasoning in LLMs. Recent systems such as OpenAI’s o1 [27] and DeepSeek-R1 [11] have shown that multi-step reasoning and self-verification can emerge purely from large-scale RL. Building on these advances, recent works [4, 8, 18, 25, 29, 38, 41, 43, 46] extend RL-enhanced reasoning to multimodal models, while SFT studies [5, 19, 41, 42, 49, 50] highlight the importance of high-quality supervision.

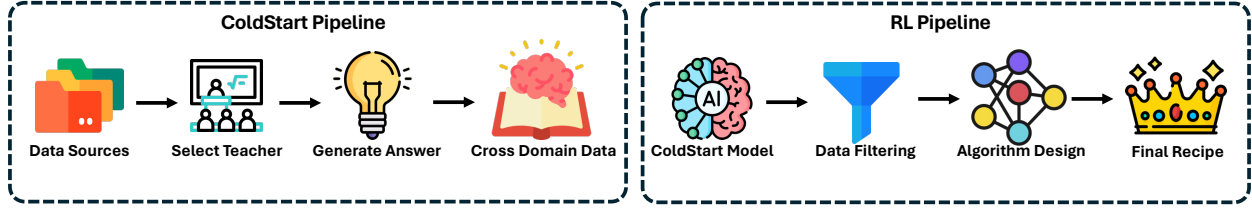


Figure 2. **Data Pipelines of OpenMMReasoner.** We propose two training recipes covering both the SFT and RL phases. The pipeline begins by collecting diverse data sources and selecting teacher models to generate new answer traces. During the RL phase, we explore different algorithm choices and filtering strategies, leading to our final optimized recipe.

Method	Data Pipeline	SFT Data	RL Data	Model
Qwen2.5-VL-7B-Instruct [3]	✗	✗	✗	✓
M2-Reasoning [2]	✗	✗	✗	✓
MiMo-VL [33]	✗	✗	✗	✓
OpenVisionReasoner [38]	✗	✓	✗	✓
OpenMMReasoner (ours)	✓	✓	✓	✓

Table 1. **Comparison of Extent of Open-Sourcing across Existing LMRMs.** OpenMMReasoner is the *first* study to fully open-source its data pipeline, SFT/RL datasets, model weights, enabling fully transparent reproduction.

SFT Datasets	Benchmarks			
Teacher Model	Average	MathVision	MathVerse	MathVista
Baseline	45.3	25.5	41.1	69.2
Qwen2.5-VL-72B-Instruct	49.8	28.1	49.7	71.5
Qwen3-VL-235B-Instruct	50.5	29.3	52.9	69.4

Table 2. **Teacher models improve performance and data efficiency.** Even with limited data, leveraging a teacher model significantly enhances the model’s reasoning ability at scale.

Despite this progress, many prior works [1, 38] lack transparent and reproducible training pipelines, and related efforts [10] focus only on textual reasoning. This limits reproducibility and obscures how data curation and training design affect outcomes. To address this, we present OpenMMReasoner, a scalable and fully open recipe covering both SFT and RL, offering practitioners clear guidance for building reliable multimodal reasoning models.

3. Supervised Fine-tuning Recipe

In this section, we present the data curation pipeline used to construct our SFT dataset. We begin with the initial data sources and describe how the dataset evolves through three stages: 1) data sourcing and formatting (103k raw questions), 2) data distillation and scaling (583k verified general-reasoning traces), and 3) cross-domain mixing to form the final 874k SFT recipe. The following experiments examine each stage and show how these design choices influence model performance. The distribution of the dataset across sources and domains is illustrated in Fig. 3.

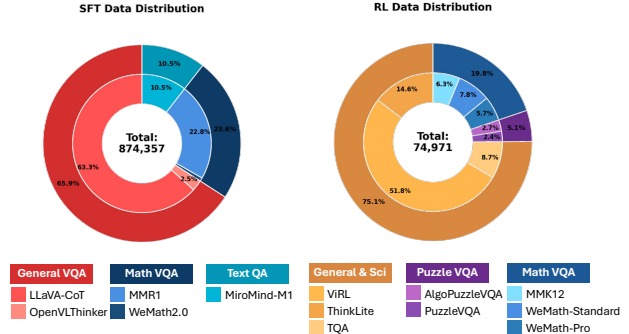


Figure 3. **Data Source Distribution OpenMMReasoner.** Our dataset comprises diverse sources across multiple domains, aiming to balance data diversity and efficiency for optimal performance.

SFT Datasets	Benchmarks			
Sampling Strategy	Average	MathVision	MathVerse	MathVista
×1 sampling	50.5	29.3	52.9	69.4
×2 sampling	52.7	30.8	54.4	72.9
×4 sampling	52.9	30.8	55.3	72.6
×8 sampling	55.2	34.6	57.1	73.7

Table 3. **Repeated teacher sampling provides a scaling axis.** With the same question sources, increasing answer diversity improves data quality, showing that both question diversity and answer diversity are important for model performance.

3.1. Data Curation

The overall pipeline is illustrated in Fig. 2. The dataset is constructed through the following stages: 1) Data Sourcing, 2) Data Distillation and Scaling, and 3) Domain Mixing.

Data Sourcing. We begin with approximately 103k raw question–answer pairs collected from public datasets including LLaVA-CoT [41], OpenVLThinker [8], and WeMath2.0 [28]. These samples mainly cover general VQA and reasoning-oriented tasks. Additional sources used later in the domain mixing stage (Section 3.5) include MMR1 [51] for image-based mathematical reasoning and MiroMind-M1 [22] for text-based mathematical reasoning.

SFT Datasets		Benchmarks		
Filter	Average	MathVision	MathVerse	MathVista
No Filter	55.2	34.6	57.1	73.7
Length Filter	54.2	33.0	56.0	73.6
Difficulty Filter	51.3	30.6	53.3	70.2

Table 4. **Over-filtering answers is detrimental.** Although filtering is commonly used in SFT, it reduces dataset diversity and does not improve overall performance.

Data Formatting. Data sourced from different benchmarks often follow inconsistent answer styles and reasoning formats, which may introduce training instability [22, 44]. We standardize all samples to a unified reasoning format, normalizing textual structure and ensuring consistent step-wise outputs before entering the distillation stage.

3.2. Experimental Setup

For the SFT training, we adopt online packing and the Liger Kernel [12] to accelerate training efficiency. All experiments are conducted using the LMMs-Engine [23] as the training framework and LMMs-Eval [20, 51] for standardized evaluation. We use Qwen2.5-VL-7B-Instruct [3] as our initial checkpoint to start, which also serves as our baseline. Each model is trained until convergence, and further implementation details are provided in the Appendix.

3.3. Data Distillation

Teacher Model Selection. Prior work [10, 11] shows that distillation from a stronger model improves data quality and data efficiency. To identify the most suitable teacher, we compare answer traces distilled from multiple candidates, applying both format validation and answer verification. Only samples whose final answers pass a rule-based validator and an LLM-as-judge check are retained, producing around 59k verified traces. As shown in Tab. 2, all teacher-based variants outperform the baseline with a clear leading, confirming that a stronger teacher yields higher-quality supervision. Both of the stronger teacher models provide an average gain of at least 4.5 points across all benchmarks, with Qwen3-VL-235B-Instruct [30] achieving the highest performance and selected as the teacher for distillation. This stage establishes the foundation for larger-scale answer sampling.

Scaling Data with verifiable answer. After selecting the optimal teacher, we enrich each question with multiple verified reasoning traces. We evaluate four scaling factors— $\times 1$, $\times 2$, $\times 4$, and $\times 8$ —and observe consistent improvements as the number of verified answers increases (Tab. 3). Quantitatively, increasing the number of verified answers per question from $\times 1$ to $\times 8$ raises the average benchmark performance from 50.5 to 55.2, demonstrating that answer diversity is a crucial factor for data quality and model generalization. We

SFT Datasets		Benchmarks		
Domain-mixing	Average	MathVision	MathVerse	MathVista
No Mix	55.2	34.6	57.1	73.7
ImgMath Mix	55.6	34.3	57.5	74.9
TxtMath Mix	55.6	35.3	57.2	74.2
Img+TxtMath	56.3	36.6	57.7	74.8

Table 5. **Cross-domain data improves reasoning ability.** Mixing data from multiple domains enhances the model’s reasoning performance, demonstrating effective reasoning transfer across domains.

adopt the $\times 8$ configuration as our final choice, as it provides substantial gains while keeping the computational cost of data generation manageable, representing a practical balance between performance and efficiency. The $\times 8$ sampling configuration increase the dataset to roughly 583k verified general-reasoning samples. This expanded dataset serves as the core of our general SFT data before domain mixing.

3.4. Length and Difficulty Filtering

We next examine whether further data refinement improves performance. Two common filtering strategies—difficulty-based filtering and length-based filtering—are evaluated. Difficulty is approximated from previous sampling accuracy, while token count is used to remove extremely short examples. As shown in Tab. 4, both filtering strategies reduce performance, which we attribute to the loss of answer diversity. Since filtering decreases sample variety without improving consistency, we adopt a no-filtering policy for the final recipe, preserving the full 583k general-reasoning set produced from the distillation stage.

3.5. Domain Mixing

To further enhance reasoning generalization, we incorporate supervised data from additional reasoning domains. While the 583k distilled dataset already provides strong multimodal reasoning, its coverage of mathematical reasoning remains limited. We therefore integrate MMR1 [18] (image-based math reasoning) and MiroMind-M1 [22] (text-based math reasoning). As shown in Tab. 5, adding both types of mathematical supervision consistently improves performance across multimodal and reasoning benchmarks. Combining the general SFT data with these additional mathematical datasets yields our final 874k mixed SFT dataset.

3.6. Analysis and Insights

In summary, the dataset evolves through three major stages—from 103k raw questions, to 583k distilled general-reasoning samples, and finally to an 874k mixed SFT dataset incorporating both general and mathematical reasoning. As shown in Tab. 6, our results across nine reasoning benchmarks [24, 28, 28, 37, 40, 47, 48, 55] show that our method achieves superior performance and data efficiency compared

Table 6. **Evaluation Results on Visual Reasoning Benchmarks.** Best results are **bold** and the second-best are underlined for *open-source models*. [†] Indicates results reproduced by ourselves.

Model	SFT Data	RL Data	MathVista testmini	MathVision test	MathVerse testmini	DynaMath worst	WeMath loose	LogicVista test	MMMU val	MMMU-Pro standard	CharXiv vision	CharXiv reas.	CharXiv desc.
<i>Close-source Models</i>													
OpenAI-GPT-4o [15]	-	-	59.9	31.1	-	40.6	34.5	-	64.4	-	-	-	-
GPT-4o mini [15]	-	-	55.1	27.3	-	30.0	31.6	48.8	41.4	-	37.6	-	34.1 74.9
<i>SFT Methods</i>													
LLaVA-OneVision-7B [21]	4.8M	-	62.6	17.6	-	17.6	9.0	-	32.0	-	24.1	-	23.6 48.7
InternVL3-8B [54]	-	-	70.5	28.6	-	33.9	23.0	-	43.6	-	-	-	37.6 73.6
Qwen2.5-VL-7B [3]	-	-	69.2	25.5	25.6 [†]	41.1	21.8	53.1	47.9	51.8 [†]	37.9 [†]	35.1 [†]	36.4 67.3
LLaVA-OneVision-1.5-8B [2]	105M	-	69.6	25.6	21.7 [†]	46.3	19.8 [†]	49.4 [†]	45.8 [†]	55.4	37.4	25.2	37.0 [†] 74.1
OMR-7B-ColdStart (ours)	874k	-	74.8	36.6	33.9	<u>57.7</u>	29.3	67.2	46.2	54.4	39.3	37.3	39.7 76.1
<i>RL-based Methods</i>													
VLLA-Thinker-Qwen2.5-7B [4]	126k	25k	68.0	26.4	-	48.2	22.4	-	48.5	-	-	-	-
ThinkLite-7B-VL [36]	-	11k	71.6	24.6	-	42.9	16.5	-	42.7	-	-	-	-
VL-Rethinker-7B [35]	-	39k	73.7	28.4	-	46.4	17.8	-	42.7	-	41.7	-	-
M2-Reasoning [1]	6.2M	102k	<u>75.0</u>	42.1	-	40.4	-	-	<u>50.6</u>	-	-	-	-
MMR1 [18]	1.6M	15k	72.0	31.8	29.0 [†]	55.4	27.9 [†]	<u>68.0[†]</u>	48.9	52.4 [†]	41.1 [†]	37.1 [†]	43.5 [†] 71.1 [†]
OpenVLThinker-7B [8]	3.3k	9.6k	65.3	23.0	26.9 [†]	38.1	16.8	61.9 [†]	44.5	<u>55.1[†]</u>	39.7 [†]	<u>38.4[†]</u>	41.0 [†] 69.2 [†]
MM-Eureka-Qwen-7B [25]	-	15.6k	72.6	28.1	32.1 [†]	45.4	23.0	59.8 [†]	46.3	54.4 [†]	40.1 [†]	37.1 [†]	42.4 [†] <u>74.1[†]</u>
OVR-7B [39]	2M	300k	72.1	51.8	<u>38.2[†]</u>	54.6	<u>33.5</u>	64.8	54.8	51.8 [†]	50.2	29.1 [†]	<u>44.5</u> 73.6
OMR-7B (ours)	874k	74k	79.5	<u>43.6</u>	38.8	63.8	34.9	79.0	50.0	57.8	<u>44.1</u>	40.6	46.1 73.5

to other SFT approaches, demonstrating the effectiveness of our recipe in building a strong reasoning base model.

Our empirical findings highlight four key observations: **(1) Answer diversity enhances reasoning.** Increasing the diversity of generated answers consistently improves the model’s overall reasoning performance, even when using the same question sources, suggesting that exposure to varied solutions strengthens understanding. **(2) Teacher model selection is crucial.** Distilling from a strong teacher model substantially boosts the model’s reasoning ability while maintaining high data efficiency. Careful selection for teacher model directly affects the quality of the distilled dataset and the final model performance. **(3) Over-filtering reduces diversity and performance.** The best results are achieved without excessive filtering, indicating that maintaining greater answer diversity encourages more robust reasoning abilities. **(4) Cross-domain knowledge improves generalization.** Incorporating diverse data from multiple domains consistently enhances the model’s overall reasoning capabilities across tasks.

4. Reinforcement Learning Recipe

To further enhance the model’s generalization and strengthen its multimodal reasoning capabilities, we introduce an RL phase tailored for multimodal reasoning tasks. Building on the strong reasoning foundation established during the SFT stage, this phase serves to further sharpen and stabilize these abilities. In this section, we provide key insights into our algorithmic design and the strategies employed to ensure stable and efficient RL dynamics.

4.1. Preliminaries of Reinforcement Learning

Let each data pair (q, a) be *i.i.d* from a distribution $\mathcal{D}$, where q is a query and a is the ground-truth answer. Given an LLM policy $\pi_\theta(\cdot|\cdot)$, let o be an LLM-generated response to q , and $r(\cdot, \cdot)$ is a predefined reward function that quantifies whether the response o yields a . RL-based fine-tuning aims to maximize this expected reward over $\mathcal{D}$, *i.e.*,

$$\max_{\theta} \mathcal{J}(\pi_\theta) \triangleq \mathbb{E}_{(q,a) \sim \mathcal{D}} \mathbb{E}_{o \sim \pi_\theta(\cdot|q)} [r(o, a)]. \quad (1)$$

Group Relative Policy Optimization (GRPO)[31] is an efficient variant of PPO that removes the need for a critic network and Generalized Advantage Estimation (GAE), thereby reducing both memory usage and computational overhead. GRPO normalizes rewards within each rollout group to reduce variance and incorporates likelihood-ratio clipping with a KL-divergence penalty to constrain π_θ close to the initial SFT policy. The GRPO objective is defined as:

$$\mathcal{J}_{\text{GRPO}}(\theta) = \mathbb{E}_{(q,a) \sim \mathcal{D}, \{o_i\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot|x)} \left[\frac{1}{G} \sum_{i=1}^G \frac{1}{|o_i|} \sum_{t=1}^{|o_i|} \left(\min \left(r_{i,t}(\theta) \hat{A}_{i,t}, \text{clip} \left(r_{i,t}(\theta), 1 - \varepsilon, 1 + \varepsilon \right) \hat{A}_{i,t} \right) \right) \right] \quad (2)$$

Decoupled Clip and Dynamic Sampling Policy Optimization (DAPO) [45] addresses several limitations of GRPO, including entropy collapse, training instability, and length bias caused by sample-level loss. To mitigate these issues, DAPO introduces a decoupled clipping mechanism and a dynamic sampling strategy, leading to a more stable and balanced optimization objective:

$$\mathcal{J}_{\text{DAPO}}(\theta) = \mathbb{E}_{(q,a) \sim \mathcal{D}, \{o_i\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot|x)} \left[\frac{1}{\sum_{i=1}^G |o_i|} \sum_{i=1}^G \sum_{t=1}^{|o_i|} \left(\min(r_{i,t}(\theta) \hat{A}_{i,t}, \text{clip}(r_{i,t}(\theta), 1 - \varepsilon_{\text{low}}, 1 + \varepsilon_{\text{high}}) \hat{A}_{i,t}) \right) \right] \quad (3)$$

subject to $0 < |\{o_j \mid \text{is_equivalent}(a, o_i)\}| < G$

Group Sequence Policy Optimization (GSPO) [53] tackles the token-level importance bias inherent in GRPO by introducing a sequence-level importance ratio for optimization. Additionally, GSPO employs a smaller clipping threshold ε to enhance training stability. The resulting optimization objective for GSPO is defined as:

$$\mathcal{J}_{\text{GSPO}}(\theta) = \mathbb{E}_{(q,a) \sim \mathcal{D}, \{o_i\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot|x)} \left[\frac{1}{G} \sum_{i=1}^G \left(\min(s_i(\theta) \hat{A}_{i,t}, \text{clip}(s_i(\theta), 1 - \varepsilon, 1 + \varepsilon) \hat{A}_{i,t}) \right) \right] \quad (4)$$

where s_i is the importance ratio based on sequence likelihood.

4.2. Dataset Curation

Dataset Sourcing Similar to the SFT stage, we begin by collecting diverse data samples across multiple domains. Our sources include MMEureka [25], ViRL [35], TQA [16], We-Math [29], PuzzleVQA [6], AlgoPuzzleVQA [9], and ThinkLiteVL [36], covering domains such as science, mathematics, charts, and puzzles.

Dataset Cleaning To ensure answer validity, we extract and verify the final answers from each dataset. We then deduplicate by computing both image and text similarities to remove redundant questions. The cleaned datasets are merged to form the final dataset, resulting in approximately 74k samples for the RL stage. Overall data statistics are shown in Fig. 3. Unless otherwise specified, all RL experiments are conducted using this dataset.

4.3. Experimental Setup

For the RL training stage, we utilize verl [32] and vllm [17] to accelerate the training process. The evaluation protocol remains consistent with that of the SFT stage, employing LMMs-Eval [20, 51] to ensure a unified and standardized evaluation setup. By default, we use temperature 1.0 and the top-performing checkpoint from the SFT phase as the starting point for RL, unless stated otherwise. The checkpoint corresponds to the best-performing model reported in Tab. 5, where cross-domain mixing of image and text-based math data were applied, along with $\times 8$ sampling for answer traces.

4.4. Training Recipe

RL Algorithm Selection To identify the most suitable RL algorithm for multimodal reasoning, we compare GSPO [53], DAPO [45], and GRPO [31] under a unified setting. We evaluate their training dynamics in terms of stability, exploration, and efficiency, as illustrated in Fig. 4. GSPO demonstrates faster convergence, higher rewards, and more stable behavior than DAPO and GRPO, achieving an effective balance between exploration and stability. While DAPO applies online filtering to remove zero-variance samples, it exhibits early entropy collapse and slower progress due to larger rollout requirements. GRPO shows moderate stability but converges more slowly. Based on these observations and the ablation results summarized in Tab. 7, we adopt GSPO as the RL algorithm in our final training setup.

Reward Function To balance task accuracy and output formatting, we adopt a composite reward function similar to those in [11, 14], which combines both aspects in a weighted manner. Specifically, the final reward R for each sample is defined as:

$$R = (1 - \lambda_{\text{fmt}}) \cdot R_{\text{acc}} + \lambda_{\text{fmt}} \cdot R_{\text{fmt}}, \quad (5)$$

where R_{acc} measures the model’s correctness on the task objective, and R_{fmt} captures the consistency of the answer format. The coefficient $\lambda_{\text{fmt}} \in [0, 1]$ controls the trade-off between task accuracy and format adherence. We use $\lambda_{\text{fmt}} = 0.1$ throughout our RL experiments.

Sampling Strategies Although RL algorithms naturally produce stronger learning signals on more challenging problems, their training efficiency remains limited. Previous work [34] employs a difficulty-based sampler, starting with easy tasks and gradually shifting to harder ones. We first use Qwen3-VL-8B [30] to track the pass rate of each problem as a difficulty metric and implement the following strategy: RL training begins with simpler tasks and progressively transitions to more complex ones. Ablation results in Tab. 7 show that curriculum learning does not outperform the original mixed sampling approach; therefore, we adopt the no-sampling strategy in our final configuration.

Balancing Reasoning Capability and Efficiency. Our empirical findings show that while OpenVisionReasoner [38] achieves competitive performance, its response length becomes excessively long, raising concerns about reasoning efficiency. This observation highlights an important question: how can a model maintain strong reasoning performance while remaining computationally efficient and adaptable across different problem types. In our reinforcement learning implementation, we adopt a similar overlength penalty strategy as proposed in DAPO [45] to mitigate the overthinking behavior and achieve a balanced trade-off between reasoning

Table 7. **Recipe selection results of different RL training strategies and coldstart starting point on Visual Reasoning Benchmarks.** Unless otherwise specified, the rollout temperature is set to 1.0 by default.

Method	Avg	MMMU	MMMU_Pro	MathVista	MathVerse	MathVision	CharXiv	LogicVista	WeMath	DynaMath			
	val	standard	vision	testmini	test	testmini	test	reas	desc	test	loose	worst	
Coldstart Start Point: $\times 8$ sampling + ImgTxt Math													
Baseline	49.4	54.4	39.3	37.3	74.8	57.7	33.9	36.6	39.7	76.1	46.2	67.2	29.3
DAPO + $\times 8$ rollout	45.6	53.3	37.7	34.6	71.7	51.0	28.6	27.9	40.7	73.0	44.4	58.3	26.4
DAPO + $\times 16$ rollout	48.9	52.9	40.6	35.6	74.3	56.2	30.3	34.5	43.4	77.0	46.9	67.6	27.4
GRPO + $\times 16$ rollout	51.1	54.6	42.8	39.4	77.1	58.3	33.9	37.1	43.1	73.8	51.1	70.2	32.1
GSPO + $\times 8$ rollout	51.6	54.9	41.2	38.8	76.9	61.4	35.9	39.9	45.0	73.8	46.9	71.7	32.3
GSPO + $\times 16$ rollout	54.3	57.8	44.1	40.6	79.5	63.8	38.8	43.6	46.1	73.5	50.0	79.0	34.9
GSPO + $\times 16$ rollout + curriculum	52.5	58.7	42.8	42.4	76.3	61.7	35.2	41.4	45.5	72.9	49.8	75.4	28.3
Coldstart Start Point: $\times 8$ sampling													
Baseline	49.2	55.6	40.6	37.7	73.7	57.1	35.2	34.6	39.5	74.1	48.9	69.1	24.0
DAPO + $\times 16$ rollout	49.6	54.8	42.0	37.6	74.1	57.8	36.8	33.9	39.1	74.3	48.7	71.6	24.0
DAPO + $\times 16$ rollout + temp. 1.4	49.3	56.9	41.0	37.3	73.3	57.8	33.2	32.2	40.1	75.0	51.6	70.2	23.4
GSPO + $\times 16$ rollout	51.0	55.1	42.4	39.6	75.1	59.4	36.2	35.4	44.6	73.0	49.1	74.0	27.9
GSPO + $\times 16$ rollout + temp. 1.4	7.4	13.4	3.9	28.3	3.4	22.0	0.0	0.0	0.2	0.0	8.3	0.2	9.2
Coldstart Start Point: $\times 1$ sampling													
Baseline	46.7	54.2	42.6	36.0	71.5	49.7	29.9	28.2	39.0	73.1	47.8	63.9	25.0
GRPO + $\times 8$ rollout	47.6	55.6	41.5	36.7	73.5	53.0	29.6	28.1	42.3	72.0	44.9	69.1	24.6
DAPO + $\times 8$ rollout	49.2	54.4	50.8	38.4	74.5	55.7	27.6	28.3	43.2	72.0	48.7	69.0	27.9
DAPO + $\times 8$ rollout + curriculum	47.0	56.0	40.5	36.8	71.6	51.2	32.9	28.1	36.2	73.6	46.2	66.8	24.6

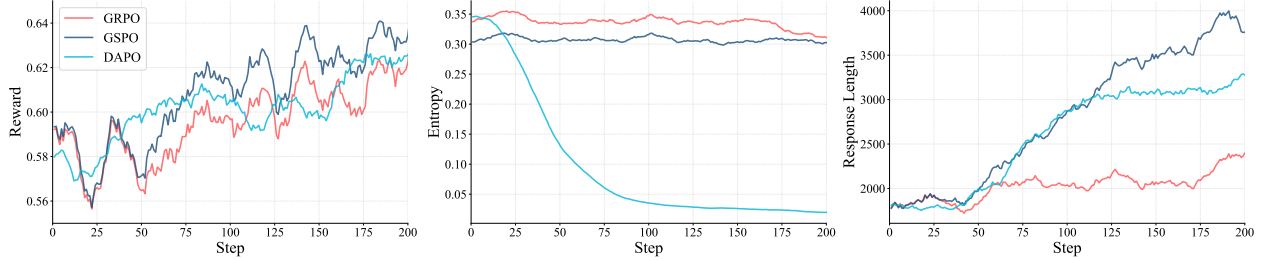


Figure 4. **Overall results across different algorithms.** We conduct a systematic comparison of various algorithms under identical multimodal RL training settings. GSPO demonstrates the highest training stability, exploration capability, and overall efficiency.

depth and efficiency. As shown in Fig. 6, we evaluate models on two benchmarks, MMMU [47] and We-Math [28]. Our model demonstrates higher reasoning efficiency compared to OVR [38]. As shown in Fig. 6, OVR requires excessively long reasoning trajectories to reach correct answers, whereas our model achieves a favorable balance between accuracy and efficiency, maintaining a reasonable reasoning budget while delivering superior overall performance.

4.5. Analysis and Insights

Overall Results. We present the evaluation results of our final models for both the SFT and RL stages in Tab. 6. Building on the strong reasoning foundation established during the SFT stage, the RL phase further sharpens and stabilizes these capabilities, leading to improved and more consistent performance. After RL, our model achieves state-of-the-art results on benchmarks such as WeMath [28], MathVerse [52], and MathVista [24] and demonstrating consistent improvement compare to SFT.

Textual Reasoning Transfers Alongside Strengthened Multimodal Reasoning. As the model’s overall reasoning ability improves through RL training, we observe the gradual emergence of textual reasoning behaviors, suggesting a transfer of reasoning competence from multimodal to purely linguistic domains. As shown in Fig. 5, validation performance on AIME24, AIME25, and AMC23 steadily increases throughout training, reflecting continuous and measurable gains in text-based reasoning capabilities. For AIME24 and AIME25, the reported accuracy represents the average over eight rollout runs. Prior work [13] has demonstrated that enhancements in mathematical reasoning can positively transfer to other reasoning domains. Evaluation results in Tab. 8 compare baseline and RL-trained models, showing that textual reasoning improves and strengthens in tandem with multimodal reasoning across all training stages. These findings extend previous observations, indicating that cross-domain reasoning skills acquired via multimodal RL can effectively transfer to purely textual tasks, further highlighting the shared cognitive foundations and underlying mechanisms across different modalities.

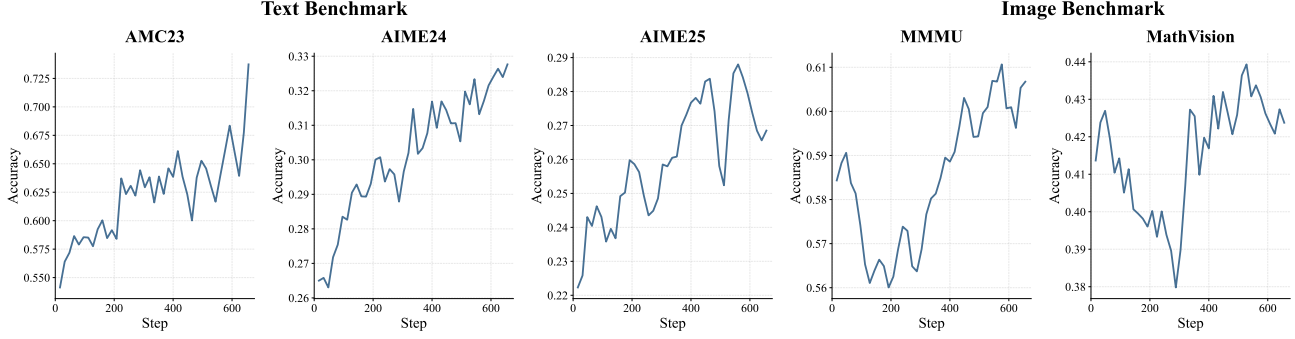


Figure 5. **Training dynamics on the validation set during RL.** During RL training, we observe that textual reasoning ability improves alongside visual reasoning, even when trained solely on multimodal data, indicating strong cross-domain generalization of reasoning capabilities.

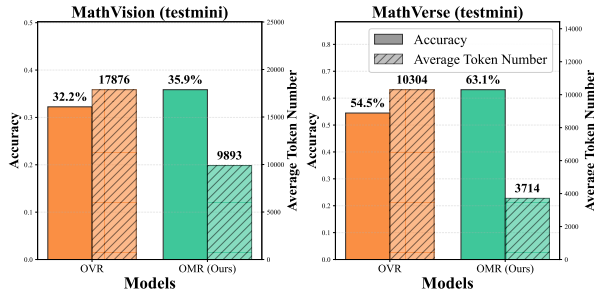


Figure 6. **Token efficiency comparison with OVR.** OpenMMReasoner achieve better accuracy while using significantly less token budget.

Benchmarks				
Model	Average	AIME24	AIME25	GPQA Diamond
Baseline	15.1	6.7	6.7	31.8
+ ColdStart	22.2	16.7	14.6	35.4
+ RL	29.4	27.1	22.1	38.9

Table 8. **Text reasoning ability transfer alongside visual reasoning improvement.** As the model’s overall reasoning capability increases, its text-based reasoning ability also improves. Across different training stages, we observe significant gains in textual reasoning, indicating strong cross-domain generalization.

Key Factors for Stable Training Dynamics. We identify two factors that critically affect RL training stability: rollout temperature and rollout count. First, higher rollout temperatures (e.g., 1.4) cause significant instability and occasional divergence, suggesting that excessive exploration amplifies policy gradient variance and destabilizes optimization. Second, the number of rollouts per update is central to maintaining convergence stability. To assess its impact, we compare $\times 8$ and $\times 16$ rollout configurations under both GRPO and DAPO. As shown in Fig. 7, the $\times 16$ configuration consistently yields higher rewards and smoother dynamics. This effect is especially pronounced for DAPO, where

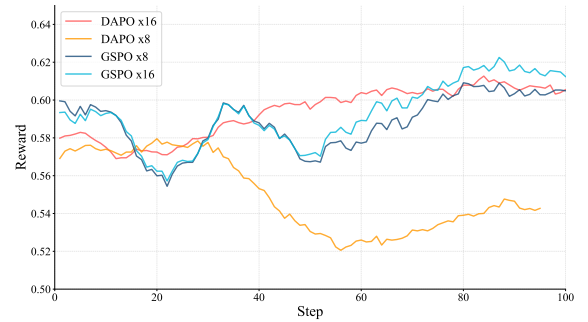


Figure 7. **Effect of rollout count on multimodal RL training stability.** DAPO becomes unstable with fewer rollouts, while increasing the rollout count leads to more stable training dynamics.

the $\times 8$ setting exhibits severe late-stage instability: entropy initially fluctuates and then abruptly spikes, ultimately causing training collapse. Interestingly, although larger rollout counts seem more computationally expensive, actual wall-clock time remains comparable between $\times 8$ and $\times 16$ due to similar token-length constraints. Based on these observations, we adopt a moderate temperature and a $\times 16$ rollout configuration as our default setup.

In summary, our key observations from the RL phase are as follows: **1) GSPO outperforms other algorithms.** GSPO demonstrates superior stability and faster convergence compared to alternative methods in multimodal RL training. **2) Token efficiency is crucial.** While increasing reasoning steps at test time can improve performance, excessive tokens reduce efficiency. Our results show that a smaller reasoning budget can achieve comparable or even better accuracy. **3) Reasoning ability transfers across domains.** Gains in reasoning during training consistently translate into stronger performance across multiple domains.

5. Conclusion

In this work, we present **OpenMMReasoner**, a transparent, scalable framework for training LMRMs via unified SFT and RL stages. Our study systematically examines how data curation, sampling strategies, and RL design shape multimodal reasoning. Experiments show that (1) well-designed SFT data, even in limited quantity, builds a strong reasoning foundation, and (2) carefully curated RL datasets, combined with effective algorithms like GSPO, enhance reasoning stability and performance. Scaling data *diversity* across domains and reasoning traces proves more valuable than merely increasing size. Structured sampling and difficulty-aware curricula improve efficiency, while well-defined reward signals strengthen reasoning precision and multimodal consistency. These insights highlight the importance of transparent, reproducible pipelines, and we hope **OpenMMReasoner** serves as a solid empirical foundation and open-source reference for scalable multimodal reasoning research.

References

- [1] Inclusion AI, Fudong Wang, Jiajia Liu, Jingdong Chen, Jun Zhou, Kaixiang Ji, Lixiang Ru, Qingpei Guo, Ruobing Zheng, Tianqi Li, et al. M2-reasoning: Empowering mllms with unified general and spatial reasoning. *arXiv preprint arXiv:2507.08306*, 2025. 3, 5
- [2] Xiang An, Yin Xie, Kaicheng Yang, Wenkang Zhang, Xiuwei Zhao, Zheng Cheng, Yirui Wang, Songcen Xu, Changrui Chen, Chunsheng Wu, et al. Llava-onevision-1.5: Fully open framework for democratized multimodal training. *arXiv preprint arXiv:2509.23661*, 2025. 3, 5
- [3] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, Humen Zhong, Yuanzhi Zhu, Mingkun Yang, Zhaohai Li, Jianqiang Wan, Pengfei Wang, Wei Ding, Zheren Fu, Yiheng Xu, Jiabo Ye, Xi Zhang, Tianbao Xie, Zesen Cheng, Hang Zhang, Zhibo Yang, Haiyang Xu, and Junyang Lin. Qwen2.5-vl technical report, 2025. 2, 3, 4, 5
- [4] Hardy Chen, Haoqin Tu, Fali Wang, Hui Liu, Xianfeng Tang, Xinya Du, Yuyin Zhou, and Cihang Xie. Sft or rl? an early investigation into training rl-like reasoning large vision-language models, 2025. 2, 5
- [5] Zesen Cheng, Sicong Leng, Hang Zhang, Yifei Xin, Xin Li, Guanzheng Chen, Yongxin Zhu, Wenqi Zhang, Ziyang Luo, Deli Zhao, and Lidong Bing. Videollama 2: Advancing spatial-temporal modeling and audio understanding in video-llms, 2024. 2
- [6] Yew Ken Chia, Vernon Toh Yan Han, Deepanway Ghosal, Lidong Bing, and Soujanya Poria. Puzzlevqa: Diagnosing multimodal reasoning challenges of language models with abstract visual patterns, 2024. 6
- [7] Gheorghe Comanici, Eric Bieber, Mike Schaekermann, Ice Pasupat, Naveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*, 2025. 2
- [8] Yihe Deng, Hritik Bansal, Fan Yin, Nanyun Peng, Wei Wang, and Kai-Wei Chang. Openvlthinker: Complex vision-language reasoning via iterative sft-rl cycles, 2025. 2, 3, 5
- [9] Deepanway Ghosal, Vernon Toh Yan Han, Chia Yew Ken, and Soujanya Poria. Are language models puzzle prodigies? algorithmic puzzles unveil serious challenges in multimodal reasoning, 2024. 6
- [10] Etash Guha, Ryan Marten, Sedrick Keh, Negin Raoof, Georgios Smyrnis, Hritik Bansal, Marianna Nezhurina, Jean Mercat, Trung Vu, Zayne Sprague, Ashima Suvana, Benjamin Feuer, Liangyu Chen, Zaid Khan, Eric Frankel, Sachin Grover, Caroline Choi, Niklas Muennighoff, Shiye Su, Wanjia Zhao, John Yang, Shreyas Pimpalgaonkar, Kartik Sharma, Charlie Cheng-Jie Ji, Yichuan Deng, Sarah Pratt, Vivek Ramanujan, Jon Saad-Falcon, Jeffrey Li, Achal Dave, Alon Albalak, Kushal Arora, Blake Wulfe, Chinmay Hegde, Greg Durrett, Sewoong Oh, Mohit Bansal, Saadia Gabriel, Aditya Grover, Kai-Wei Chang, Vaishaal Shankar, Aaron Gokaslan, Mike A. Merrill, Tatsunori Hashimoto, Yejin Choi, Jenia Jitsev, Reinhard Heckel, Maheswaran Sathiamoorthy, Alexandros G. Dimakis, and Ludwig Schmidt. Openthoughts: Data recipes for reasoning models, 2025. 2, 3, 4
- [11] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025. 2, 4, 6
- [12] Pin-Lun Hsu, Yun Dai, Vignesh Kothapalli, Qingquan Song, Shao Tang, Siyu Zhu, Steven Shimizu, Shivam Sahni, Haowen Ning, Yanning Chen, and Zhipeng Wang. Liger-kernel: Efficient triton kernels for LLM training. In *Championing Open-source DEvelopment in ML Workshop @ ICML25*, 2025. 4
- [13] Maggie Huan, Yuetai Li, Tuney Zheng, Xiaoyu Xu, Seung-gone Kim, Minxin Du, Radha Poovendran, Graham Neubig, and Xiang Yue. Does math reasoning improve general llm capabilities? understanding transferability of llm reasoning, 2025. 7
- [14] Hugging Face. Open r1: A fully open reproduction of deepseek-r1, 2025. 6
- [15] Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, et al. Gpt-4o system card. *arXiv preprint arXiv:2410.21276*, 2024. 5
- [16] Aniruddha Kembhavi, Min Joon Seo, Dustin Schwenk, Jonghyun Choi, Ali Farhadi, and Hannaneh Hajishirzi. Are you smarter than a sixth grader? textbook question answering for multimodal machine comprehension. In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 5376–5384. IEEE Computer Society, 2017. 6
- [17] Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the ACM SIGOPS 29th Symposium on Operating Systems Principles*, 2023. 6, 1
- [18] Sicong Leng, Jing Wang, Jiayi Li, Hao Zhang, Zhiqiang Hu, Boqiang Zhang, Yuming Jiang, Hang Zhang, Xin Li, Lidong Bing, Deli Zhao, Wei Lu, Yu Rong, Aixin Sun, and Shijian Lu. Mmr1: Enhancing multimodal reasoning with variance-aware sampling and open resources, 2025. 2, 4, 5
- [19] Ang Li, Charles Wang, Deqing Fu, Kaiyu Yue, Zikui Cai, Wang Bill Zhu, Ollie Liu, Peng Guo, Willie Neiswanger, Furong Huang, Tom Goldstein, and Micah Goldblum. Zebra-cot: A dataset for interleaved vision language reasoning, 2025. 2
- [20] Bo Li*, Peiyuan Zhang*, Kaichen Zhang*, Fanyu Pu*, Xinrun Du, Yuhao Dong, Haotian Liu, Yuanhan Zhang, Ge Zhang, Chunyuan Li, and Ziwei Liu. Lmms-eval: Accelerating the development of large multimodal models, 2024. 4, 6
- [21] 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. 5
- [22] Xingxuan Li, Yao Xiao, Dianwen Ng, Hai Ye, Yue Deng, Xiang Lin, Bin Wang, Zhanfeng Mo, Chong Zhang, Yueyi

- Zhang, Zonglin Yang, Ruilin Li, Lei Lei, Shihao Xu, Han Zhao, Weiling Chen, Feng Ji, and Lidong Bing. Miromind-m1: An open-source advancement in mathematical reasoning via context-aware multi-stage policy optimization, 2025. 3, 4
- [23] LMMs-Lab. Lmms engine: A simple, unified multimodal framework for pretraining and finetuning., 2025. 4
- [24] 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, 2024. 4, 7
- [25] Fanqing Meng, Lingxiao Du, Zongkai Liu, Zhixiang Zhou, Quanfeng Lu, Daocheng Fu, Tiancheng Han, Botian Shi, Wenhai Wang, Junjun He, et al. Mm-eureka: Exploring the frontiers of multimodal reasoning with rule-based reinforcement learning. *arXiv preprint arXiv:2503.07365*, 2025. 2, 5, 6
- [26] OpenAI. Introducing o3 and o4-mini. <https://openai.com/index/introducing-o3-and-o4-mini/>, 2025. 2
- [27] OpenAI. Learning to reason with llms. <https://openai.com/index/learning-to-reason-with-llms/>, 2025. 2
- [28] 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. 3, 4, 7
- [29] Runqi Qiao, Qiuna Tan, Peiqing Yang, Yanzi Wang, Xiaowan Wang, Enhui Wan, Sitong Zhou, Guanting Dong, Yuchen Zeng, Yida Xu, Jie Wang, Chong Sun, Chen Li, and Honggang Zhang. We-math 2.0: A versatile mathbook system for incentivizing visual mathematical reasoning, 2025. 2, 6
- [30] Qwen Team. Qwen3-vl: Sharper vision, deeper thought, broader action. <https://qwen.ai/blog?id=99f0335c4ad9ff6153e517418d48535ab6d8afef&from=research.latest-advancements-list>, 2025. Accessed: 2025-11-14. 4, 6
- [31] 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. 2, 5, 6
- [32] Guangming Sheng, Chi Zhang, Zilingfeng Ye, Xibin Wu, Wang Zhang, Ru Zhang, Yanghua Peng, Haibin Lin, and Chuan Wu. Hybridflow: A flexible and efficient rlhf framework. *arXiv preprint arXiv: 2409.19256*, 2024. 6
- [33] Core Team, Zihao Yue, Zhenru Lin, Yifan Song, Weikun Wang, Shuhuai Ren, Shuhao Gu, Shicheng Li, Peidian Li, Liang Zhao, Lei Li, Kainan Bao, Hao Tian, Hailin Zhang, Gang Wang, Dawei Zhu, Cici, Chenhong He, Bowen Ye, Bowen Shen, Zihan Zhang, Zihan Jiang, Zhixian Zheng, Zhichao Song, Zhenbo Luo, Yue Yu, Yudong Wang, Yuanyuan Tian, Yu Tu, Yihan Yan, Yi Huang, Xu Wang, Xinzhe Xu, Xingchen Song, Xing Zhang, Xing Yong, Xin Zhang, Xiangwei Deng, Wenyu Yang, Wenhan Ma, Weiwei Lv, Weiwei Zhuang, Wei Liu, Sirui Deng, Shuo Liu, Shima Chen, Shihua Yu, Shaohui Liu, Shande Wang, Rui Ma, Qiantong Wang, Peng Wang, Nuo Chen, Menghang Zhu, Kangyang Zhou, Kang Zhou, Kai Fang, Jun Shi, Jinhao Dong, Jiebao Xiao, Jiaming Xu, Huaqiu Liu, Hongshen Xu, Heng Qu, Haochen Zhao, Hanglong Lv, Guoan Wang, Duo Zhang, Dong Zhang, Di Zhang, Chong Ma, Chang Liu, Can Cai, and Bingquan Xia. Mimo-vl technical report, 2025. 3
- [34] Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, Chuning Tang, Congcong Wang, Dehao Zhang, Enming Yuan, Enzhe Lu, Fengxiang Tang, Flood Sung, Guangda Wei, Guokun Lai, Haiqing Guo, Han Zhu, Hao Ding, Hao Hu, Hao Yang, Hao Zhang, Haotian Yao, Haotian Zhao, Haoyu Lu, Haoze Li, Haozhen Yu, Hongcheng Gao, Huabin Zheng, Huan Yuan, Jia Chen, Jianhang Guo, Jianlin Su, Jianzhou Wang, Jie Zhao, Jin Zhang, Jingyuan Liu, Junjie Yan, Junyan Wu, Lidong Shi, Ling Ye, Longhui Yu, Mengnan Dong, Neo Zhang, Ningchen Ma, Qiwei Pan, Qucheng Gong, Shaowei Liu, Shengling Ma, Shupeng Wei, Sihan Cao, Siying Huang, Tao Jiang, Weihao Gao, Weimin Xiong, Weiran He, Weixiao Huang, Weixin Xu, Wenhao Wu, Wenyang He, Xianghui Wei, Xianqing Jia, Xingzhe Wu, Xinran Xu, Xinxing Zu, Xinyu Zhou, Xuehai Pan, Y. Charles, Yang Li, Yangyang Hu, Yangyang Liu, Yanru Chen, Yejie Wang, Yibo Liu, Yidao Qin, Yifeng Liu, Ying Yang, Yiping Bao, Yulun Du, Yuxin Wu, Yuzhi Wang, Zaida Zhou, Zhaoji Wang, Zhaowei Li, Zhen Zhu, Zheng Zhang, Zhexu Wang, Zhilin Yang, Zhiqi Huang, Zihao Huang, Ziyao Xu, Zonghan Yang, and Zongyu Lin. Kimi k1.5: Scaling reinforcement learning with llms, 2025. 6
- [35] Haozhe Wang, Chao Qu, Zuming Huang, Wei Chu, Fangzhen Lin, and Wenhui Chen. Vl-rethinker: Incentivizing self-reflection of vision-language models with reinforcement learning, 2025. 5, 6
- [36] Xiyao Wang, Zhengyuan Yang, Chao Feng, Hongjin Lu, Linjie Li, Chung-Ching Lin, Kevin Lin, Furong Huang, and Lijuan Wang. Sota with less: Mcts-guided sample selection for data-efficient visual reasoning self-improvement, 2025. 2, 5, 6
- [37] Zirui Wang, Mengzhou Xia, Luxi He, Howard Chen, Yitao Liu, Richard Zhu, Kaiqu Liang, Xindi Wu, Haotian Liu, Sadhika Malladi, Alexis Chevalier, Sanjeev Arora, and Danqi Chen. Charxiv: Charting gaps in realistic chart understanding in multimodal llms, 2024. 4
- [38] Yana Wei, Liang Zhao, Jianjian Sun, Kangheng Lin, Jisheng Yin, Jingcheng Hu, Yinmin Zhang, En Yu, Haoran Lv, Zejia Weng, Jia Wang, Chunrui Han, Yang Peng, Qi Han, Zheng Ge, Xiangyu Zhang, Daxin Jiang, and Vishal M. Patel. Open vision reasoner: Transferring linguistic cognitive behavior for visual reasoning, 2025. 2, 3, 6, 7
- [39] Yana Wei, Liang Zhao, Jianjian Sun, Kangheng Lin, Jisheng Yin, Jingcheng Hu, Yinmin Zhang, En Yu, Haoran Lv, Zejia Weng, et al. Open vision reasoner: Transferring linguistic cognitive behavior for visual reasoning. *arXiv preprint arXiv:2507.05255*, 2025. 5
- [40] Yijia Xiao, Edward Sun, Tianyu Liu, and Wei Wang. Log-

- icvista: Multimodal llm logical reasoning benchmark in visual contexts, 2024. 4
- [41] Guowei Xu, Peng Jin, Ziang Wu, Hao Li, Yibing Song, Lichao Sun, and Li Yuan. Llava-cot: Let vision language models reason step-by-step, 2025. 2, 3
- [42] Yi Yang, Xiaoxuan He, Hongkun Pan, Xiyan Jiang, Yan Deng, Xingtao Yang, Haoyu Lu, Dacheng Yin, Fengyun Rao, Minfeng Zhu, Bo Zhang, and Wei Chen. R1-onevision: Advancing generalized multimodal reasoning through cross-modal formalization, 2025. 2
- [43] Zuhao Yang, Sudong Wang, Kaichen Zhang, Keming Wu, Sicong Leng, Yifan Zhang, Bo Li, Chengwei Qin, Shijian Lu, Xingxuan Li, and Lidong Bing. Longvt: Incentivizing "thinking with long videos" via native tool calling. *arXiv preprint arXiv:2511.20785*, 2025. 2
- [44] Yixin Ye, Zhen Huang, Yang Xiao, Ethan Chern, Shijie Xia, and Pengfei Liu. Limo: Less is more for reasoning, 2025. 4
- [45] Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Weinan Dai, Tiantian Fan, Gaohong Liu, Lingjun Liu, Xin Liu, Haibin Lin, Zhiqi Lin, Bole Ma, Guangming Sheng, Yuxuan Tong, Chi Zhang, Mofan Zhang, Wang Zhang, Hang Zhu, Jinhua Zhu, Jiaze Chen, Jiangjie Chen, Chengyi Wang, Hongli Yu, Yuxuan Song, Xiangpeng Wei, Hao Zhou, Jingjing Liu, Wei-Ying Ma, Ya-Qin Zhang, Lin Yan, Mu Qiao, Yonghui Wu, and Mingxuan Wang. Dapo: An open-source llm reinforcement learning system at scale, 2025. 2, 5, 6
- [46] Ruifeng Yuan, Chenghao Xiao, Sicong Leng, Jianyu Wang, Long Li, Weiwen Xu, Hou Pong Chan, Deli Zhao, Tingyang Xu, Zhongyu Wei, Hao Zhang, and Yu Rong. VI-cogito: Progressive curriculum reinforcement learning for advanced multimodal reasoning, 2025. 2
- [47] Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruofei 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. 4, 7
- [48] Xiang Yue, Tianyu Zheng, Yuansheng Ni, Yubo Wang, Kai Zhang, Shengbang Tong, Yuxuan Sun, Botao Yu, Ge Zhang, Huan Sun, Yu Su, Wenhui Chen, and Graham Neubig. Mmmu-pro: A more robust multi-discipline multimodal understanding benchmark, 2025. 4
- [49] Boqiang Zhang, Kehan Li, Zesen Cheng, Zhiqiang Hu, Yuqian Yuan, Guanzheng Chen, Sicong Leng, Yuming Jiang, Hang Zhang, Xin Li, Peng Jin, Wenqi Zhang, Fan Wang, Lidong Bing, and Deli Zhao. Videollama 3: Frontier multimodal foundation models for image and video understanding, 2025. 2
- [50] Hang Zhang, Xin Li, and Lidong Bing. Video-llama: An instruction-tuned audio-visual language model for video understanding, 2023. 2
- [51] Kaichen Zhang, Bo Li, Peiyuan Zhang, Fanyi Pu, Joshua Adrian Cahyono, Kairui Hu, Shuai Liu, Yuanhan Zhang, Jingkang Yang, Chunyuan Li, and Ziwei Liu. Lmms-eval: Reality check on the evaluation of large multimodal models, 2024. 3, 4, 6
- [52] Renrui Zhang, Dongzhi Jiang, Yichi Zhang, Haokun Lin, Ziyu Guo, Pengshuo Qiu, Aojun Zhou, Pan Lu, Kai-Wei Chang, Peng Gao, and Hongsheng Li. Mathverse: Does your multimodal llm truly see the diagrams in visual math problems?, 2024. 7
- [53] Chuji Zheng, Shixuan Liu, Mingze Li, Xiong-Hui Chen, Bowen Yu, Chang Gao, Kai Dang, Yuqiong Liu, Rui Men, An Yang, Jingren Zhou, and Junyang Lin. Group sequence policy optimization, 2025. 2, 6
- [54] Jinguo Zhu, Weiyun Wang, Zhe Chen, Zhaoyang Liu, Shenglong Ye, Lixin Gu, Hao Tian, Yuchen Duan, Weijie Su, Jie Shao, Zhangwei Gao, Erfei Cui, Xuehui Wang, Yue Cao, Yangzhou Liu, Xingguang Wei, Hongjie Zhang, Haomin Wang, Weiye Xu, Hao Li, Jiahao Wang, Nianchen Deng, Songze Li, Yanan He, Tan Jiang, Jiapeng Luo, Yi Wang, Conghui He, Botian Shi, Xingcheng Zhang, Wenqi Shao, Junjun He, Yingdong Xiong, Wenwen Qu, Peng Sun, Penglong Jiao, Han Lv, Lijun Wu, Kaipeng Zhang, Huipeng Deng, Jiaye Ge, Kai Chen, Limin Wang, Min Dou, Lewei Lu, Xizhou Zhu, Tong Lu, Dahua Lin, Yu Qiao, Jifeng Dai, and Wenhui Wang. Internvl3: Exploring advanced training and test-time recipes for open-source multimodal models, 2025. 5
- [55] Chengke Zou, Xingang Guo, Rui Yang, Junyu Zhang, Bin Hu, and Huan Zhang. Dynamath: A dynamic visual benchmark for evaluating mathematical reasoning robustness of vision language models, 2025. 4

OpenMMReasoner: Pushing the Frontiers for Multimodal Reasoning with an Open and General Recipe

Supplementary Material

Training Component	SFT	RL
optimizer	AdamW	AdamW
scheduler	cosine	constant
learning rate	5e-5	1e-6
weight decay	0.0	0.1
Training Steps	4300	1232
Warmup Steps	430	25
Max Length	61440	32792
Dynamic Bsz	True	True
Remove padding	True	True
Liger Kernel	True	False

Table 1. Detail parameters for SFT and RL training.

1. Implementation Details

1.1. Training Details

We describe the training procedures for our best-performing models in both the SFT and RL stages.

SFT. To maximize training throughput and reduce memory consumption, we apply online stream packing with an iterable dataset, removing padding tokens to avoid unnecessary computation. We use a packing length of 61,440 tokens. For each batch, we dynamically compute the input token lengths of incoming samples and fill the batch buffer until it is fully packed. Because packing is performed online, the total number of epochs cannot be predetermined; instead, we train each model until convergence. The full set of experimental hyperparameters is provided in Tab. 1.

RL. For RL training, we use a global batch size of 128, which strikes a balance between on-policy stability and training speed. Log probabilities are computed using a dynamic batch-size implementation without padding to further reduce memory usage. During generation, we set the maximum number of new tokens to 28,696 and cap the prompt length at 4,096 tokens, resulting a max length at 32792. We use a temperature of 1.0 and keep this configuration fixed across all experiments. Due to the high computational cost of RL training, we run training until the reward saturates. The detailed hyperparameters are listed in Tab. 1.

1.2. Evaluation Details

We now describe our evaluation setup for multimodal reasoning benchmarks. The SFT and RL models share the same

evaluation configuration. We use the system prompt shown in Tab. 4 to ensure the model outputs both the reasoning trace and the final answer in an extractable format. The extracted answer is then validated using a two-stage process: a rule-based validator followed by an LLM-as-judge validator. We first apply the rule-based validator to minimize evaluation cost; if the answer cannot be verified, we fall back to the LLM-as-judge, using the prompt provided in Tab. 5. For all evaluations, we use a temperature of 0.0 for reproducibility and set the maximum generation length to 49K tokens, except for AIME, where we use a temperature of 1.0. We employ vLLM [17] as the serving engine to accelerate inference.

2. Additional Result and Analysis

Sampling Scaling Results. We present the full evaluation results for different sampling strategies in Tab. 2. As shown in the table, scaling up the sampling strategy improves the average score from 46.7 to 49.2, demonstrating the effectiveness of increasing answer diversity along this axis.

Rollout Analysis. During RL training, we record all rollout logs and analyze the proportion of reflection-related words in the model’s outputs. We observe that as the reward increases over training steps, the proportion of reflection words also rises. To illustrate this, we generate a word cloud from the final-step rollout outputs—after removing noise words—and find that reflection cues such as “let,” “wait,” and “think” appear frequently in the model’s responses. This indicates that our RL recipe effectively encourages the model to engage in more explicit reasoning as its capabilities improve. The results are shown in Fig. 1.

Reward λ_{fmt} Ablation. To assess the impact of the λ_{fmt} parameter in RL, we conduct additional experiments using four values: 0.1, 0.3, 0.5, and 0.7, under the same experimental setup as the GRPO configuration. The evaluation results are presented in Tab. 3. We observe that a lower format-reward weight, specifically $\lambda_{fmt} = 0.1$, consistently yields the best performance. Consequently, we adopt $\lambda_{fmt} = 0.1$ in our final configuration.

3. Examples

Data examples. We present several examples in Tab. 6 and Tab. 7 from our reasoning dataset to demonstrate the high quality of the answer traces generated by our method.

Table 3. Reward ablation result.

Method	Avg	MMMU	MMMU_Pro	MathVista	MathVerse	MathVision	CharXiv	LogicVista	WeMath	DynaMath			
	val	standard	vision	testmini	test	testmini	test	reas	desc	test	loose	worst	
Reward λ_{fmt} Settings													
$\lambda_{fmt} = 0.1$	51.1	54.6	42.8	39.4	77.1	58.3	33.9	37.1	43.1	73.8	51.1	70.2	32.1
$\lambda_{fmt} = 0.3$	47.0	51.3	36.7	33.9	73.9	55.8	29.6	34.1	36.5	75.9	43.3	64.3	29.1
$\lambda_{fmt} = 0.5$	45.0	48.3	34.3	32.4	74.2	53.1	26.0	27.9	39.8	75.6	41.1	58.1	29.3
$\lambda_{fmt} = 0.7$	48.8	53.4	37.8	35.8	74.8	56.8	32.9	35.8	39.4	75.7	46.9	66.9	29.1

System Prompt for model

You are a helpful assistant. When the user asks a question, your response must include two parts: first, the reasoning process
 $\rightarrow$ enclosed in <think>...</think> tags, then the final answer enclosed in <answer>...</answer> tags. Please provide a
 $\rightarrow$ clear, concise response within <answer> </answer> tags that directly addresses the question.

Table 4. The prompt that used in evaluation.

System Prompt for judge model

You are a strict evaluator assessing answer correctness. You must output 1 for fully correct answers and 0 for any other case.
 $\rightarrow$ You will receive the question, the ground truth answer, and the model prediction.

Input

Question:

...

{question}

...

Ground Truth Answer:

...

{answer}

...

Model Prediction:

...

{prediction}

...

Evaluation Rules

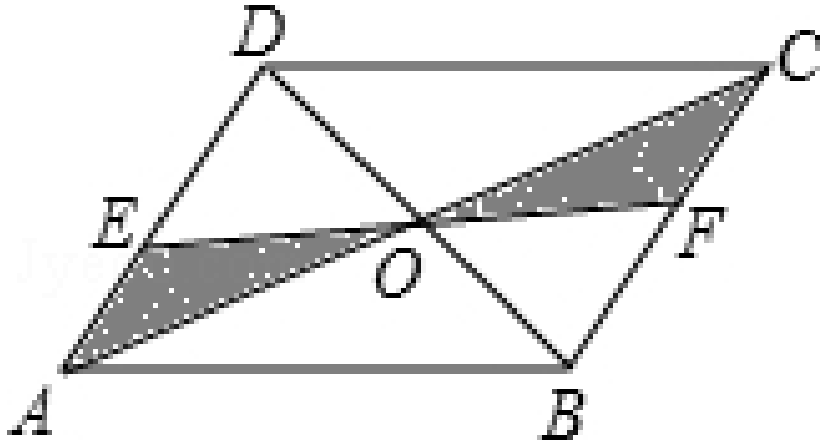
- For multiple-choice questions: Score 1 if the predicted answer matches the ground truth answer, it can be directly in option
 $\rightarrow$ letters or the content of the options.
- For open-ended questions:
 - * Score 1 if the prediction matches the answer semantically, it can be in different format.
 - * Score 0 for partially correct answers or answers with extra incorrect information, even if the reasoning process is correct.
- Ignore minor differences in formatting, capitalization, or spacing since the model may explain in a different way.
- Treat numerical answers as correct if they match within reasonable precision
- For questions requiring units, both value and unit must be correct

Strict Output format

1 or 0

Table 5. The prompt that used in evaluation.

ColdStart Data Samples



Question As shown in the figure, line EF passes through the intersection point O of the diagonals of the parallelogram and intersects AD and BC at points E and F, respectively. If the area of the parallelogram is 12, the sum of the areas of triangles AOE and BOF equals () A. 2 B. 3 C. 4 D. Cannot be determined Output the thinking process in <think> </think> and final answer in <answer> </answer> tags.

Answer <think>In a parallelogram, the diagonals bisect each other, so point O is the midpoint of both diagonals AC and BD. Also, since EF passes through O and intersects AD and BC at E and F respectively, we can use symmetry and properties of parallelograms. Because O is the center of the parallelogram, any line through O divides the parallelogram into two regions of equal area. Moreover, triangles AOE and COF are congruent (by ASA or SAS, since $AO = OC$, $\angle AOE = \angle COF$, and $\angle OAE = \angle OCF$ due to parallel lines and transversal), and similarly, triangles BOF and DOE are congruent. Therefore, the area of triangle AOE equals the area of triangle COF, and the area of triangle BOF equals the area of triangle DOE. The total area of the parallelogram is 12. The diagonals divide it into 4 triangles of equal area: AOB, BOC, COD, DOA — each of area 3. Now, consider the two shaded triangles: AOE and BOF. Since EF passes through O, and by symmetry, the sum of the areas of AOE and BOF is equal to the area of triangle AOB (which is one of the four equal-area triangles formed by the diagonals). Why? Because: $\text{Area}(\text{AOE}) + \text{Area}(\text{BOF}) = \text{Area}(\text{AOB})$. This is because the parts not included in AOE and BOF (namely, EOD and FOC) are congruent to AOE and BOF respectively, so swapping them doesn't change the total. Alternatively, since the entire parallelogram is divided into 4 triangles of area 3 each, and $\text{AOE} + \text{BOF}$ together make up exactly one of those triangles (AOB), their combined area is 3. </think><answer>B. 3</answer>

Table 6. An example of the coldstart data sample.

ColdStart Data Samples

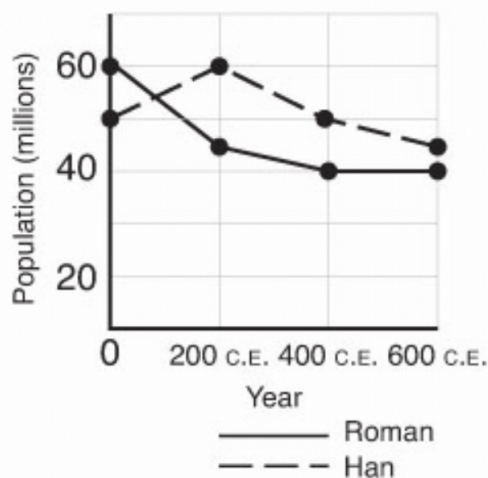


Question What is the man doing with the beans?A. cookingB. eating themC. countingD. selling themAnswer with the option's letter from the given choices directly. Output the thinking process in <think> </think> and final answer in <answer> </answer> tags.

Answer <think>The image shows a man sitting beside a cart loaded with green beans in what appears to be a busy street market. He is handling the beans and seems to be interacting with potential customers, which is typical behavior for a vendor. There is no visible cooking equipment, nor is he eating or counting the beans. The context strongly suggests he is selling them.</think><answer>D</answer>

Table 7. An example of the coldstart data sample.

No. 25: Which of the following best explains the overall trend shown in the [image 1] ?



Options:

- | | |
|--|---|
| (A) Political instability leading to population decline | (B) The spread of pathogens across the Silk Road |
| (C) Development of new trade routes | (D) Climate change affecting the Silk Road |
| (E) Migrations to areas of Central Asia for resettlement | (F) Technological advancements in transportation |
| (G) Invasions by Mongol tribes | (H) Large-scale famine due to crop failures |
| (I) Economic prosperity and population growth | (J) Rise of religious conflicts along the Silk Road |

Previous Next

1 / 170 Go

Question

A. Political instability leading to population decline B. The spread of pathogens across the Silk Road C. Development of new trade routes D. Climate change affecting the Silk Road E. Migrations to areas of Central Asia for resettlement F. Technological advancements in transportation G. Invasions by Mongol tribes H. Large-scale famine due to crop failures I. Economic prosperity and population growth J. Rise of religious conflicts along the Silk Road Answer the question in the image.

Answer

B

Ours

<think>The graph shows population trends for Roman and Han populations over time, with a general decline in population for both over centuries. This decline is most plausibly explained by widespread disease, as the Silk Road facilitated the spread of pathogens across regions, leading to population reductions. While political instability or other factors could contribute, the most direct and historically supported cause for such a widespread population decline across empires is the spread of pathogens.</think><answer>B. The spread of pathogens across the Silk Road</answer>

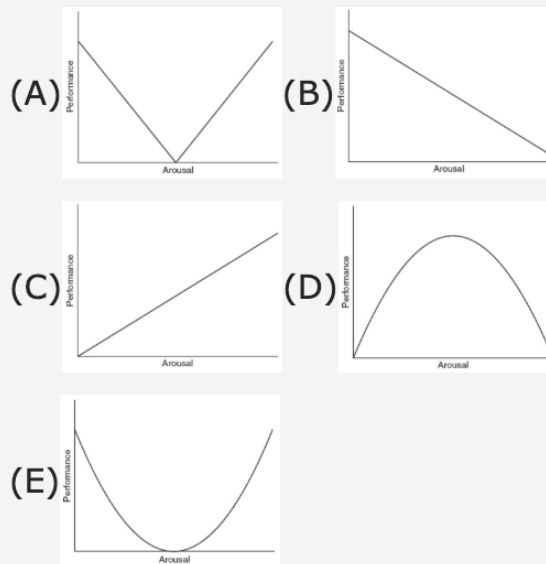
Qwen2.5-VL-7B-Instruct

<think>The graph shows a general decline in population for both the Roman and Han empires over the period from 200 C.E. to 600 C.E. This trend is not consistent with economic prosperity, technological advancements, or large-scale famine, which would typically lead to population growth. The spread of pathogens across the Silk Road could potentially cause a decline in population, but the graph does not provide specific information about the cause of the decline. The rise of religious conflicts along the Silk Road could also lead to population decline, but the graph does not provide specific information about the cause of the decline. The development of new trade routes and the rise of religious conflicts along the Silk Road are not directly related to the population decline shown in the graph. The graph does not provide specific information about the cause of the decline, but it is reasonable to infer that the decline in population could be due to political instability, which could lead to a decline in the overall population.</think><answer>A. Political instability leading to population decline</answer>

Table 8. An example of our model inference result

No. 10: Which of the following graphs best displays the relationship between a person's level of arousal and his or her performance on a challenging task?

Options:



Question	A. <image 3> B. <image 2> C. <image 4> D. <image 5> E. <image 1> Answer the question in the image.
Answer	D
Ours	<think> The relationship between arousal and performance on a challenging task is typically described by the Yerkes-Dodson Law, which states that performance increases with arousal up to a point, after which further increases in arousal lead to a decline in performance. This forms a bell-shaped curve 2014 peak performance at moderate arousal, with performance decreasing as arousal rises too high or falls too low. Looking at the options: - (A) Inverted V: performance improves with low arousal, then drops 2014 doesn't match. - (B) Straight line decreasing: performance always decreases with arousal 2014 too linear and wrong direction. - (C) Straight line increasing: performance always increases with arousal 2014 incorrect. - (D) Bell curve: performance peaks at moderate arousal 2014 matches Yerkes-Dodson Law. - (E) U-shape: performance improves with low arousal, then drops 2014 opposite of what's expected. Thus, (D) best represents the relationship. The question asks to choose from A-E, and (D) is labeled as image 5 in the options mapping. So, the correct answer is D. </think><answer>D. <image 5></answer>
Qwen2.5-VL-7B-Instruct	<think>The relationship between a person's level of arousal and their performance on a challenging task is often described by the Yerkes-Dodson Law, which suggests that performance increases with arousal up to a certain point and then decreases. This relationship is typically represented by a bell-shaped curve, indicating an optimal level of arousal for peak performance.</think><answer>E</answer>

Table 9. An example of our model inference result