

A Survey of Mathematical Reasoning in the Era of Multimodal Large Language Model: Benchmark, Method & Challenges

Anonymous ACL submission

Abstract

Mathematical reasoning, a core aspect of human cognition, is vital across many domains, from educational problem-solving to scientific advancements. As artificial general intelligence (AGI) progresses, integrating large language models (LLMs) with mathematical reasoning tasks is becoming increasingly significant. This survey provides the **first comprehensive analysis of mathematical reasoning in the era of multimodal large language models (MLLMs)**. We review over 200 studies published since 2021, and examine the state-of-the-art developments in Math-LLMs, with a focus on multimodal settings. We categorize the field into three dimensions: *benchmarks*, *methodologies*, and *challenges*. In particular, we explore multimodal mathematical reasoning pipeline, as well as the role of (M)LLMs and the associated methodologies. Finally, we identify seven major challenges hindering the realization of AGI in this domain, offering insights into the future direction for enhancing multimodal reasoning capabilities. This survey serves as a critical resource for the research community in advancing the capabilities of LLMs to tackle complex multimodal reasoning tasks.

1 Introduction

Mathematical reasoning is a critical aspect of human cognitive ability, involving the process of deriving conclusions from a set of premises through logical and systematic thinking (Jonsson et al., 2022; Yu et al., 2024b). It plays an essential role in a wide range of applications, from problem-solving in education to advanced scientific discoveries. As artificial general intelligence (AGI) continues to advance (Zhong et al., 2024), the integration of large language models (LLMs) with mathematical reasoning tasks becomes increasingly significant. These models, with their impressive capabilities in language understanding, have the potential to simulate complex reasoning processes that were once

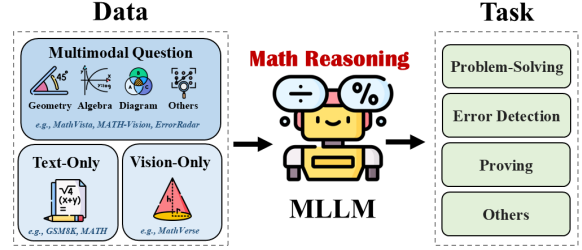


Figure 1: The illustration of our research scope (*i.e.*, investigating the MLLM’s math reasoning capability).

Survey	Venue & Year	Scope	Multimodal	LLM
(O’Halloran, 2015)	JMB’15	MM4Math	✓	
(Hegedus and Tall, 2015)	IRME’15	MM4Math	✓	
(Lu et al., 2022b)	ACL’22	DL4Math		✓
(Li et al., 2023a)	arXiv’23	LLM4Edu		✓
(Liu et al., 2023b)	arXiv’23	LLM4Edu		✓
(Li et al., 2024g)	COLM’24	DL4TP		✓
(Ahn et al., 2024)	EACL’24	LLM4Math		✓
(Xu et al., 2024a)	IJMLC’24	LLM4Edu		✓
(Wang et al., 2024d)	arXiv’24	LLM4Edu		✓
Ours	-	MLLM4Math	✓	✓

Table 1: Comparisons between relevant surveys & ours.

thought to be inherently human. In recent years, both academia and industry have placed increasing emphasis on this direction (Wang et al., 2024d; Xu et al., 2024a; Lu et al., 2022b; Yan et al., 2025).

The inputs for mathematical reasoning tasks are diverse, extending beyond traditional text-only to multimodal settings, as illustrated in Figure 1. Mathematical problems often involve not only textual information but also visual elements, such as diagrams, graphs, or equations, which provide essential context for solving the problem (Wang et al., 2024e; Yin et al., 2024). In the past year, multimodal mathematical reasoning has emerged as a key focus for multimodal large language models (MLLMs) (Zhang et al., 2024c; Bai et al., 2024; Wu et al., 2023a). This shift is driven by the recognition that reasoning tasks in fields like mathematics require models capable of integrating and processing multiple modalities simultaneously to achieve human-like performance. However, multimodal mathematical reasoning poses significant

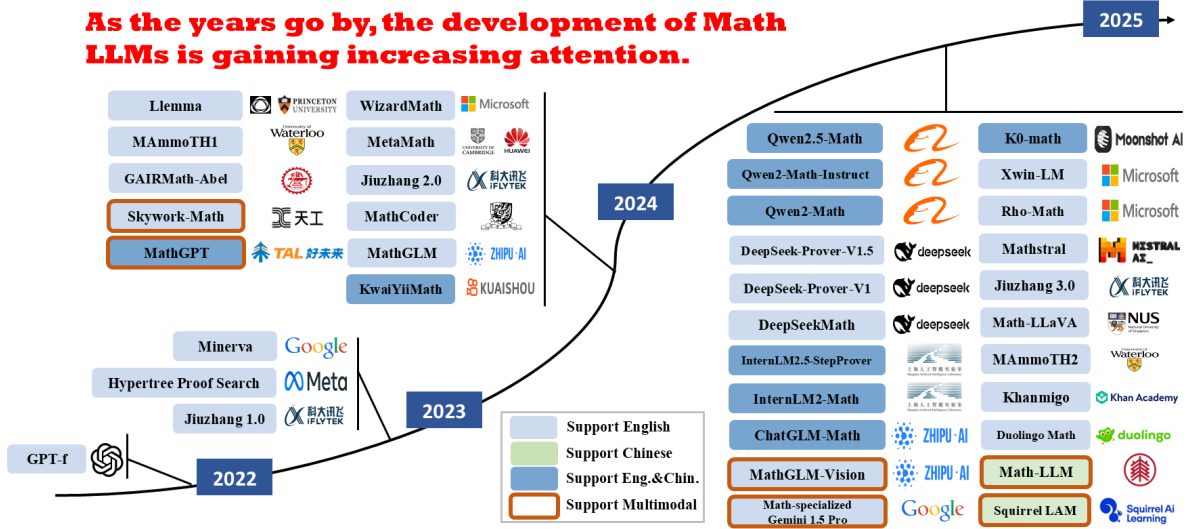


Figure 2: The release timeline of Math-LLMs in recent years.

challenges due to the complex interaction between different modalities, the need for deep semantic understanding, and the importance of context preservation across modalities (Liang et al., 2024a; Song et al., 2023; Fu et al., 2024b). These challenges are central to the realization of AGI, where models must integrate diverse forms of knowledge seamlessly to perform sophisticated reasoning tasks.

Math-LLM Progress. Figure 2 illustrates that, driven by the rapid development of LLMs since 2021, the number of math-specific LLMs (Math-LLMs) has grown steadily, alongside enhanced support for multilingual and multimodal capabilities (More details in Appendix A). The landscape was marked by the introduction of models like GPT-f (Polu and Sutskever, 2021) and Minerva (Lewkowycz et al., 2022), with Hypertree Proof Search (Lample et al., 2022) and Jiuzhang 1.0 (Zhao et al., 2022) highlighting advancements in theorem proving and mathematical question understanding capabilities, respectively. Year 2023 saw a surge in diversity and specialization, alongside multimodal support from models like Skywork-Math (Zeng et al., 2024). In year 2024, there was a clear focus on enhancing mathematical instruction (e.g., Qwen2.5-Math (Yang et al., 2024a)) and proof (e.g., DeepSeek-Prover (Xin et al., 2024a)) capabilities. The year also witnessed the emergence of Math-LLMs with a vision component, such as MathGLM-Vision (Yang et al., 2024b).

Scope. Previous surveys have not fully captured the progress and challenges of mathematical reasoning in the age of MLLMs. As indicated in Table

1, some works have concentrated on the application of deep learning techniques to mathematical reasoning (Lu et al., 2022b) or specific domains such as theorem proving (Li et al., 2024g), but they have overlooked the rapid advancements brought about by the rise of LLMs. Others have broadened the scope to include the role of LLMs in education (Wang et al., 2024d; Xu et al., 2024a; Li et al., 2023a) or mathematical fields (Ahn et al., 2024; Liu et al., 2023b), but have failed to explore the development and challenges of mathematical reasoning in multimodal settings in depth. Therefore, this survey aims to fill this gap by providing the **first-ever comprehensive analysis of the current state of mathematical reasoning in the era of MLLMs**, focusing on three key dimensions: *benchmark*, *methodology*, and *challenges*.

Structure. In this paper, we survey over 200 publications from the AI community since 2021 related to (M)LLM-based mathematical reasoning, and summarize the progress of Math-LLMs. We first approach the field from the benchmark perspective, analyzing the LLM-based mathematical reasoning task through four key aspects: basic focus, task, evaluation, and training data (Section 2). Subsequently, we explore the roles that (M)LLMs play in mathematical reasoning, categorizing them as reasoners, enhancers, and planners (Section 3). Finally, we identify seven core challenges that the mathematical reasoning faces in the era of MLLMs (Section 4). This survey aims to provide the community with comprehensive insights for advancing multimodal reasoning capabilities of LLMs.

2 Benchmark Perspective

2.1 Overview

Benchmarking for mathematical reasoning plays a crucial role in advancing LLM research, as it provides standardized, reproducible pipeline for assessing the performance on reasoning tasks. While previous benchmarks such as GSM8K (Cobbe et al., 2021) and MathQA (Amini et al., 2019) were instrumental in the pre-LLM era, our scope is centered on those relevant to (M)LLMs. In this section, we present a comprehensive analysis of recent benchmarks for mathematical reasoning in the context of (M)LLMs (Shown in Table 3 from Appendix B). The section is organized into four subsections: Basic Focus (Sec.2.2), Tasks (Sec.2.3), Evaluation (Sec.2.4), and Training Data (Sec.2.5).

2.2 Basic Focus

Basic Format. In a math reasoning task (taking problem-solving as a basic setting), the goal is to solve a mathematical problem given a specific format of input and output. The input consists of a statement that describes the problem to be solved. As shown in Figure 3, this can be presented in either a textual format or a multimodal format (text accompanied by visual elements, such as figures or diagrams). The output is the predicted solution to the problem, represented as numerical or symbolic results. More cases can be seen in Appendix C.

Language & Size. The majority of benchmarks are available in English, with a few exceptions like Chinese (Li et al., 2024i) or Romanian (Cosma et al., 2024) datasets. This predominance of English datasets underscores the challenges of multilingual representation in the mathematical reasoning domain, suggesting an opportunity for future work to diversify datasets across languages, especially those in underrepresented regions. Moreover, the size of these datasets varies widely, from smaller sets (e.g., QRData (Liu et al., 2024d) with 411 questions) to massive corpora (e.g., OpenMathInstruct-1 (Toshniwal et al., 2024b) with 1.8 million problem-solution pairs). Larger datasets are more likely to support robust model training and evaluation, but their size can also present challenges in terms of computational requirements and quality control.

Source. The sources of datasets predominantly consist of public (i.e., derived from public repositories or datasets) and private sources. The private datasets typically offer specialized problem types

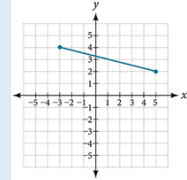
(a) Text-only Math Reasoning Setting

[Qns] Find the distance between the two endpoints using the distance formula. The two endpoints of the line are $(-3, 4)$ and $(5, 2)$, respectively.

[Ans] 8.246

(b) Multimodal Math Reasoning Setting

[Qns] Find the distance between the two endpoints.



[Ans] 8.246

Figure 3: Typical data format of math reasoning task for text-only & multimodal settings. Examples are derived from MathVerse (Zhang et al., 2024f), which assess whether and how much MLLMs can truly understand the visual diagrams for mathematical reasoning.

and tasks, and may present unique challenges, such as restricted access or ethical considerations. On the other hand, public datasets foster wider community collaboration, though they may suffer from limitations in diversity and task coverage. Some works have also leveraged LLMs to generate the datasets tailored to specific needs. For instance, GeomVerse constructs synthetic datasets to evaluate the multi-hop reasoning abilities required in geometric math problems (Kazemi et al., 2023).

Educational Level. The benchmarks span various educational levels, ranging from elementary school to university-level problems. Besides, there has also been a surge in datasets focused on competition-level problems (Tsoukalas et al.), offering insights into the current limitations of LLMs in comparison to the upper bound of human cognitive abilities. Future directions could involve more focused datasets targeting specific educational levels to enable models to specialize in handling particular age groups or skill sets.

2.3 Task

Model Choice. The choice of models in these benchmarks spans open-source and closed-source models, with a growing interest in Math-LLMs. This trend indicates an increasing recognition of the need for models tailored to mathematical reasoning, which often require specialized training and handling of structured knowledge. Additionally, with the recent release of GPT-4o (OpenAI, 2024)

and Gemini-Pro-1.5 (Reid et al., 2024), which have demonstrated significant advancements in multimodal reasoning capabilities, the latest benchmarks have begun to include them in the evaluations. For example, ErrorRadar, in its initial formulation of multimodal error detection setting, incorporates these state-of-the-art MLLMs to highlight the real-world performance gap between AI systems and human-level reasoning (Yan et al., 2024a).

Reasoning Task. Problem-solving tasks typically dominate, reflecting the emphasis on students’ ability to apply knowledge and reasoning skills in real-world contexts. This also serves as the core objective of current Math-LLMs. In addition, a growing proportion of error detection tasks suggests an increasing focus on helping students recognize and correct mistakes (Li et al., 2024e; Yan et al., 2024a; Kurtic et al., 2024). Meanwhile, proving tasks, often associated with higher-order thinking, highlight a shift towards cultivating logical reasoning and systematic problem-solving abilities (Tsoukalas et al.). Moreover, a smaller portion of work has addressed tasks that align with real-world educational needs but lack systematic formulation. For instance, Li et al. (2024e) further introduces error correction (which goes beyond simple error detection); Didolkar et al. (2024) explores automated skill discovery for problem-solving; and MathChat (Liang et al., 2024c) focuses on reasoning in multi-turn settings (such as follow-up QA and problem generation). Given the higher demands on reasoning capabilities in multimodal settings, many studies have also evaluated the aforementioned reasoning tasks in image-text problem settings. These efforts aim to provide the LLM community with more diverse, real-world task scenarios, catering to the needs of multimodal learning environments.

2.4 Evaluation

Discriminative Evaluation is a common approach, focusing on the ability of M(LLM)s to correctly classify or choose the correct answer (Hendrycks et al., 2021; Mishra et al., 2022; Li et al., 2024c). Based on specific motivations, some works also build their metrics upon accuracy for further expansion. For example, GSM-PLUS, a new adversarial benchmark for evaluating the robustness of LLMs in mathematical reasoning, develops performance drop rate (PDR) to measure the relative decline in performance on question variations compared to the original questions (Li et al., 2024d). ErrorRadar uses error step accuracy and error category

accuracy together to evaluate the multimodal error detection of MLLMs (Yan et al., 2024a).

Generative Evaluation, on the other hand, measures a M(LLM)’s ability to produce detailed explanations or solve problems from scratch. This evaluation type is gaining traction, particularly for complex mathematical tasks where step-by-step solutions are required. For instance, MathVerse, which modifies problems with varying degrees of information content in multi-modality, employs GPT-4 to score each key step in the reasoning process generated by MLLMs (Zhang et al., 2024f). CHAMP proposes a solution evaluation pipeline where GPT-4 is utilized as a grader for the answer summary, given the ground truth answer (Mao et al., 2024).

Due to page limit, more details of both types of evaluation metrics can be seen in Appendix D.

2.5 Training Data

The training of MLLMs for mathematical reasoning relies on a carefully orchestrated integration of *instruction design*, *data scale*, and *task diversity* to ensure robust and generalizable performance. Central to this process is the **design of instruction sets**, which are structured to bridge symbolic, textual, and visual reasoning (Toshniwal et al., 2024b,a). These instructions progressively escalate in complexity, starting from foundational arithmetic to advanced domains like calculus and linear algebra, ensuring models build skills incrementally. Each problem can be accompanied by explicit step-by-step explanations, enabling models to learn logical sequencing and self-correction (Zhang et al., 2024g; Tang et al., 2024b; Liang et al., 2024c).

The **scale of pre-training data** also plays an equally critical role. Models are exposed to terabytes of data sourced from textbooks, research papers (e.g., arXiv), online educational platforms (e.g., Khan Academy), and synthetically generated problems. A significant portion (10–30%) of the pretraining corpus is dedicated to mathematical content, with specialized datasets ensuring coverage of niche topics. While scaling to trillion-token corpora enhances robustness, rigorous filtering mechanisms, such as self-supervised quality checks, are applied to eliminate noise, including incorrect solutions or irrelevant content (Shao et al., 2024; Qwen, 2024; Yue et al., 2024c).

Finally, the **variety of mathematical tasks** ensures models adapt to diverse challenges. Training spans core domains like algebra and geometry, as well as cross-disciplinary applications (e.g.,

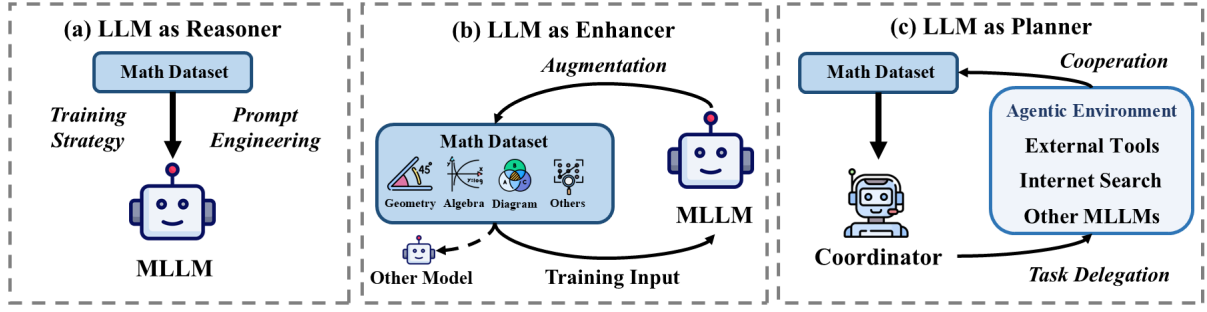


Figure 4: The illustration of the comparisons among three paradigms of (M)LLM-based mathematical reasoning.

physics-based calculus problems). Tasks are presented in multiple formats: closed-ended questions (e.g., solving equations), open-ended prompts (e.g., deriving proofs), and error-analysis exercises that require identifying and correcting flawed reasoning (Lu et al., 2022b; Yan et al., 2025).

For example, G-LLaVA (Gao et al., 2023) focuses on solving geometry problems by extracting visual features from geometric figures and jointly modeling them with text descriptions, allowing the model to understand key elements (e.g., points, lines, angles) in geometric figures and their relationship with text descriptions. MAVIS (Zhang et al., 2024g) features an automatic data generation engine that can quickly generate large-scale, high-quality multimodal mathematical datasets, addressing the problem of data scarcity. It also uses instruction fine-tuning to teach the model how to decompose complex mathematical problems and generate reasonable reasoning steps (esp., MAVIS-Instruct includes 834k visual math problems with CoT rationales). Math-LLaVA (Shi et al., 2024) uses the MathV360K multimodal dataset (360k instances), which covers multiple mathematical domains to gradually improve the model’s mathematical reasoning ability through bootstrapping and further optimize the model using generated data.

3 Methodology Perspective

3.1 Overview & Findings

MLLMs have been leveraged in various ways to tackle the broad spectrum of mathematical reasoning tasks. Based on our comprehensive review of recent methodologies (summarized in Table 5 from Appendix E), we classify the works into three distinct paradigms: LLM as Reasoner (Sec.3.2), LLM as Enhancer (Sec.3.3), and LLM as Planner (Sec.3.4), and finally provide a in-depth comparison of technical distinctions (Sec.3.5).

Findings. First, single-modality settings dominate the current landscape of method-oriented research, with the majority focusing solely on algebraic tasks. However, since 2024, multimodal approaches have been increasingly incorporated, expanding the scope of mathematical reasoning to include geometry, diagrams, and even broader mathematical concepts. This shift signals a growing interest in enhancing model robustness through multimodal learning, which can address the diverse nature of mathematical problems. Second, regarding the evaluated tasks, problem-solving and proving are gaining prominence, while some research also focuses on error detection or others (e.g., RefAug includes error correction and follow-up QA as evaluation tasks (Zhang et al., 2024j)). Finally, in terms of the role of LLMs, Reasoner is the most common role, followed by Enhancer, while Planner remains less explored but holds promise due to recent advancements in multi-agent intelligence.

3.2 LLM as Reasoner

Definition. In the *Reasoner* paradigm, M(LLM)s harness their inherent reasoning capabilities to solve mathematical problems, as shown in Figure 4 (a). This can either involve fine-tuning existing LLMs on task-specific datasets or utilizing zero-shot or few-shot learning strategies. These models utilize advanced semantic understanding and reasoning techniques, such as symbolic manipulation, logical deduction, and multi-step reasoning.

Examples. Deng et al. (2023) develops a unified framework for answer calibration that integrates step-level and path-level strategies on multi-step reasoning of LLMs. MATH-SHEPHERD serves as a process-oriented math verifier, which assigns a reward score to each step of the LLM’s outputs on math questions (Wang et al., 2024c). As for multimodal approaches, Math-PUMA introduces progressive upward multimodal alignment strat-

egy for reasoning-enhanced training (Zhuang et al., 2024); Math-LLaVA, a LLaVA-1.5-based model, directly bootstraps mathematical reasoning via fine-tuned on 360K high-quality math QA pairs, which can ensure the depth and breadth of multimodal mathematical problems (Shi et al., 2024); STIC develops a two-stage self-training pipeline (consisting of Image Comprehension Self-Training phase & Description-Infused Fine-Tuning phase) for enhancing visual comprehension (Deng et al., 2024); VCAR emphasizes on the visual-centric supervision, thus proposing a similar two-step training pipeline which handles the visual description generation task first, followed by mathematical rationale generation task (Jia et al., 2024).

Summary & Outlook. This paradigm has shown significant promise, particularly in solving problems requiring multiple steps of reasoning. However, despite improvements, issues with robustness remain, particularly with zero-shot reasoning tasks. Future work should focus on combining reasoning with structured knowledge retrieval systems and enhancing models’ ability to reason effectively across diverse domains, especially in multimodal contexts (Fan et al., 2024b; Pan et al., 2023).

3.3 LLM as Enhancer

Definition. In the *Enhancer* paradigm, M(LLM)s are primarily used to augment data, thereby enabling improvements in mathematical reasoning, as illustrated in Figure 4 (b). This can be achieved by synthesizing new training data, refining existing datasets, or introducing new variations that target specific problem-solving abilities (Li et al., 2022). Data augmentation can include paraphrasing mathematical problems, adding noise to mathematical expressions, or generating problem variants for underrepresented cases.

Examples. A typical example of a single-modality enhancement approach is Masked Thought, which introduces perturbations to the input and randomly masks tokens within the chain of thought during training (Chen et al., 2024a). Math-Genie, which aims to generate diverse and reliable math problems and solution from a small-scale dataset, leverages a solution augmentation model to iteratively create new solutions from existing ones (Lu et al., 2024b). For multimodal methods, AlphaGeometry proves most olympiad-level mathematical theorems, via trained from scratch on large-scale synthetic data guiding the symbolic deduction (Trinh et al., 2024); LogicSolver intro-

duces interpretable formula-based tree-structure for each solution equation (Yang et al., 2022); InfiMM-Math achieves the exceptional performance as it is trained on a large-scale multimodal interleaved math dataset developed and validated by LLMs such as LLaMA3-70B-Instruct (Han et al., 2024); DFE-GPS constructs its synthetic training set, which integrates visual features and geometric formal language (Zhang et al., 2024i).

Summary & Outlook. This paradigm offers substantial performance improvements by enriching the training set. However, challenges remain in ensuring the diversity and relevance of the generated data. Moreover, while text-based augmentation methods have proven effective, the potential for multimodal augmentation is still underexplored. Future research should focus on advancing multimodal data augmentation techniques, especially for tasks that require interaction between visual and textual modalities (Xiao et al., 2023).

3.4 LLM as Planner

Definition. In the *Planner* paradigm, M(LLM)s are treated as coordinators that guide the solution of complex mathematical problems by delegating tasks to other models or tools, as illustrated in Figure 4 (c). This includes scenarios where multiple agents or models collaborate to achieve a single objective, thereby enhancing the performance of mathematical problem-solving through cooperative interactions. These models often work in environments with multiple steps or require iterative refinement of solutions.

Examples. A notable tool-integrated agent is ToRA, which plans the sequential use of natural language rationale and program-based tools synergistically to solve mathematical problems in an optimal manner (Gou et al., 2023). Additionally, COPRA simulates a single agent-like reasoning mechanism where GPT-4 proposes tactic applications within a stateful backtracking search, leveraging feedback from the proof environment (Thakur et al., 2024). This can also extend to multimodal scenarios, as seen in Chameleon, which serves as an AI system that augments MLLMs with plug-and-play modules for compositional reasoning, leveraging an LLM-based planner to assemble tools for complex tasks (Lu et al., 2024a). Furthermore, Visual Sketchpad presents the concept of sketching as a ubiquitous tool used by humans for communication, ideation, and problem-solving. Hence, MLLMs can enable external tools (e.g., matplotlib)

Aspect	LLM as Reasoner	LLM as Enhancer	LLM as Planner
Data Interaction Patterns			
<i>Input-Output Relation</i>	End-to-end mapping (Problem $\rightarrow$ Answer)	Data augmentation pipeline (Raw data $\rightarrow$ Enhanced data)	Dynamic workflow planning (Problem $\rightarrow$ Plan $\rightarrow$ Subtasks)
<i>External Dependencies</i>	Low (Self-contained reasoning)	Medium (Data distribution dependent)	High (Requires toolchain integration)
Pros & Cons			
<i>Advantages</i>	Transparent reasoning & Strong interpretability	Improves generalization & Handles data scarcity	Breaks capability boundaries & Enables complex task solving
<i>Limitations</i>	Error-prone in complex reasoning	May introduce semantic biases	High system complexity & Increased latency

Table 2: Comparisons among the three methodology paradigms.

to generate intermediate sketches to aid in reasoning, which includes an iterative interaction process with an environment (Hu et al., 2024). Although there has been much work on Compositional Visual Reasoning in the past (Gupta and Kembhavi, 2023; Surís et al., 2023; Yao et al., 2022), Visual Sketchpad is the first work that integrates the planning capabilities of MLLMs with the real gap of mathematical reasoning settings (*i.e.*, sketch-based reasoning involving visuo-spatial concepts).

Summary & Outlook. While the Planner paradigm introduces significant improvements, particularly for complex tasks that require multi-agent collaboration, it remains a relatively under-explored area (Xi et al., 2023; Guo et al., 2024b). There is potential for further improvement in task decomposition, agent cooperation strategies, and integration of diverse computational tools. Future work will likely focus on refining these planning strategies, especially for multimodal systems that can jointly leverage visual and textual knowledge to solve more intricate problems (Xie et al., 2024; Durante et al., 2024; Li et al., 2023b).

3.5 Paradigm Comparison

As summarized in Table 2, we list the differences between the three paradigms to provide the community with a more comprehensive understanding of the latest technical distinctions. These three paradigms show a progressive development logic: Reasoner focuses on intrinsic model capabilities, Enhancer targets data optimization, and Planner moves towards system-level intelligent collaboration. In practice, we also anticipate adopting a hybrid approach (*e.g.*, using Enhancer to generate augmented data to train Reasoner, then coordinating multiple Reasoner modules via Planner to solve complex problems). This layered architecture may become the core design paradigm for future multimodal mathematical reasoning systems.

4 Challenges

In the realm of MLLMs for mathematical reasoning, the following key challenges persist that hinder their full potential. Addressing these challenges is essential for advancing MLLMs toward more robust and flexible systems that can better support mathematical reasoning in real-world settings.

❶ **Lack of High-Quality, Diverse, and Large-Scale Multimodal Datasets.** As discussed in Section 2.5, current multimodal mathematical reasoning datasets face tripartite limitations in quality (*e.g.*, misaligned text-image pairs), scale (insufficient advanced topic coverage), and task diversity (overemphasis on problem-solving versus error diagnosis or theorem proving). For instance, most datasets focus on question answering but lack annotations for error tracing steps or formal proof generation, while synthetic datasets often exhibit domain bias (Wang et al., 2024a; Lu et al., 2023). Three concrete solutions emerge: i) Develop hybrid dataset construction pipelines combining expert-curated problems with AI-augmented task variations; ii) Implement cross-task knowledge distillation, where models trained on proof generation guide error diagnosis through attention pattern transfer; iii) Leverage automated frameworks quality-controlled multimodal expansion to systematically generate diverse task formats (*e.g.*, converting proof exercises into visual dialogues). More discussion on data bottlenecks in Appendix F.1.

❷ **Insufficient Visual Reasoning.** Many math problems require extracting and reasoning over visual content, such as charts, tables, or geometric diagrams. Current models struggle with intricate visual details, such as interpreting three-dimensional geometry or analyzing irregularly structured tables (Zhang et al., 2024f). Hence, it may be beneficial to introduce enhanced visual feature extraction modules and integrate scene graph representations for better reasoning over complex visual elements (Ibrahim et al., 2024; Guo et al., 2024c).

③ Reasoning Beyond Text and Vision. While the current research focus on the combination of text and vision, mathematical reasoning in real-world applications often extends beyond these two modalities. For instance, audio explanations, interactive problem-solving environments, or dynamic simulations might play a role in some tasks. Current models are not well-equipped to handle such diverse inputs (Abrahamson et al., 2020; Jusslin et al., 2022). To address this, datasets should be expanded to include more diverse modalities, such as audio, video, and interactive tools. MLLMs should also be designed with flexible architectures capable of processing and reasoning over multiple types of inputs, allowing for a richer representation of mathematical problems (Dasgupta et al., 2023).

④ Limited Domain Generalization. Mathematical reasoning spans many domains, such as algebra, geometry, diagram and commonsense, each with its own specific requirements for problem-solving (Liu et al., 2023b; Lu et al., 2022b). Math-LLMs that perform well in one domain often fail to generalize across others, which can limit their utility. By pretraining and fine-tuning Math-LLMs on a wide array of problem types, models may handle cross-domain tasks more effectively, improving their ability to generalize across different mathematical topics and problem-solving strategies. We extend more discussion on limited domain generalization in multimodal contexts in Appendix F.2.

⑤ Error Feedback Limitations. Mathematical reasoning involves various types of errors, such as calculation mistakes, logical inconsistencies, and misinterpretations of the problem. Currently, MLLMs lack mechanisms to detect, categorize, and correct these errors effectively, which can result in compounding mistakes throughout the reasoning process (Yan et al., 2024a; Li et al., 2024e). A potential solution is to integrate error detection and classification modules that can identify errors at each step of the reasoning process. Besides, multi-agent collaboration mechanism could be introduced, via involving multiple agents collaborating by exchanging feedback and collectively refining the reasoning process (Xu et al., 2024d). We extend more discussion on error feedback limitation in multimodal contexts in Appendix F.3.

⑥ Integration with Real-world Educational Needs. Existing benchmarks and models often overlook real-world educational contexts, such as how students use draft work, like handwritten notes or diagrams, to solve problems (Xu et al., 2024c;

Wang et al., 2024d). These real-world elements are crucial for understanding how humans approach mathematical reasoning (Mouchere et al., 2011; Gervais et al., 2024). By incorporating draft notes, handwritten calculations, and dynamic problem-solving workflows into the training data, MLLMs can be tailored to provide more accurate and contextually relevant feedback for students.

⑦ Test-Time Scaling Technique in Multimodal Context. While foundation models increasingly adopt test-time scaling techniques (*e.g.*, dynamic architecture adaptation), their integration with multimodal mathematical reasoning remains underexplored and suboptimal. For example, current implementations like o1 (Jaech et al., 2024) or DeepSeek-R1 (Guo et al., 2025) struggle to dynamically allocate computational resources based on math problem complexity across modalities, such as deciding when to prioritize symbolic computation over visual parsing for optimization problems. Future work should focus on two directions: i) Develop modality-aware scaling controllers that jointly consider problem type, visual complexity, and required mathematical operations to optimize dynamic architecture decisions; ii) Create lightweight meta-optimization layers that can adjust model capacity allocation (*e.g.*, expert selection in MoE systems) through real-time analysis of multimodal problem-solving workflows (Xu et al., 2025a; Besta et al., 2025). Such advancements could enable more efficient trade-offs between accuracy and computational cost in deployed systems. We also discuss how test-time scaling techniques can tackle the other challenges in Appendix F.4.

5 Conclusion

In this survey, we have provided a comprehensive overview of the progress and challenges in mathematical reasoning within the context of MLLMs. We highlighted the significant advances in the development of Math-LLMs and the growing importance of multimodal integration for solving complex reasoning tasks. We identified five key challenges that are crucial for the continued development of AGI systems capable of performing sophisticated mathematical reasoning tasks. As research continues to advance, it is essential to focus on these challenges to unlock the full potential of LLMs in multimodal settings. We hope this survey provides insights to guide future LLM research, ultimately leading to more effective and human-like mathematical reasoning capabilities in AI systems.

Limitations

Despite our best efforts to ensure comprehensive coverage of the published works, it is possible that some relevant studies were overlooked. Additionally, human errors could have occurred during the categorization or referencing of papers in the survey. To minimize such errors, we made a concerted effort to gather studies from multiple sources and performed a multiple-round checking process. While minor inconsistencies or omissions may still exist, we believe this survey represents the most comprehensive review of MLLM-based mathematical reasoning to date, effectively capturing key research trends and highlighting ongoing challenges.

References

Dor Abrahamson, Mitchell J Nathan, Caro Williams-Pierce, Candace Walkington, Erin R Ottmar, Hortensia Soto, and Martha W Alibali. 2020. The future of embodied design for mathematics teaching and learning. In *Frontiers in Education*, volume 5, page 147. Frontiers Media SA.

Janice Ahn, Rishu Verma, Renze Lou, Di Liu, Rui Zhang, and Wenpeng Yin. 2024. Large language models for mathematical reasoning: Progresses and challenges. *arXiv preprint arXiv:2402.00157*.

Aida Amini, Saadia Gabriel, Peter Lin, Rik Koncel-Kedziorski, Yejin Choi, and Hannaneh Hajishirzi. 2019. Mathqa: Towards interpretable math word problem solving with operation-based formalisms. *arXiv preprint arXiv:1905.13319*.

Zhangir Azerbayev, Hailey Schoelkopf, Keiran Paster, Marco Dos Santos, Stephen McAleer, Albert Q Jiang, Jia Deng, Stella Biderman, and Sean Welleck. 2023. Llemma: An open language model for mathematics. *arXiv preprint arXiv:2310.10631*.

Tianyi Bai, Hao Liang, Binwang Wan, Yanran Xu, Xi Li, Shiyu Li, Ling Yang, Bozhou Li, Yifan Wang, Bin Cui, et al. 2024. A survey of multimodal large language model from a data-centric perspective. *arXiv preprint arXiv:2405.16640*.

Maciej Besta, Julia Barth, Eric Schreiber, Ales Kubicek, Afonso Catarino, Robert Gerstenberger, Piotr Nyczyk, Patrick Iff, Yueling Li, Sam Houlston, et al. 2025. Reasoning language models: A blueprint. *arXiv preprint arXiv:2501.11223*.

Changyu Chen, Xiting Wang, Ting-En Lin, Ang Lv, Yuchuan Wu, Xin Gao, Ji-Rong Wen, Rui Yan, and Yongbin Li. 2024a. Masked thought: Simply masking partial reasoning steps can improve mathematical reasoning learning of language models. *arXiv preprint arXiv:2403.02178*.

Feng Chen, Allan Raventos, Nan Cheng, Surya Ganguli, and Shaul Druckmann. 2025. Rethinking fine-tuning when scaling test-time compute: Limiting confidence improves mathematical reasoning. *arXiv preprint arXiv:2502.07154*.

Jiaqi Chen, Tong Li, Jinghui Qin, Pan Lu, Liang Lin, Chongyu Chen, and Xiaodan Liang. 2022. Unigeo: Unifying geometry logical reasoning via reformulating mathematical expression. *arXiv preprint arXiv:2212.02746*.

Jiaqi Chen, Jianheng Tang, Jinghui Qin, Xiaodan Liang, Lingbo Liu, Eric P Xing, and Liang Lin. 2021. Geoqa: A geometric question answering benchmark towards multimodal numerical reasoning. *arXiv preprint arXiv:2105.14517*.

Qiguang Chen, Libo Qin, Jin Zhang, Zhi Chen, Xiao Xu, and Wanxiang Che. 2024b. M3cot: A novel benchmark for multi-domain multi-step multi-modal chain-of-thought. *arXiv preprint arXiv:2405.16473*.

Wenhu Chen, Ming Yin, Max Ku, Pan Lu, Yixin Wan, Xueguang Ma, Jianyu Xu, Xinyi Wang, and Tony Xia. 2023. Theoremqa: A theorem-driven question answering dataset. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 7889–7901.

Zhe Chen, Weiyun Wang, Yue Cao, Yangzhou Liu, Zhangwei Gao, Erfei Cui, Jinguo Zhu, Shenglong Ye, Hao Tian, Zhaoyang Liu, et al. 2024c. Expanding performance boundaries of open-source multimodal models with model, data, and test-time scaling. *arXiv preprint arXiv:2412.05271*.

Anoop Cherian, Kuan-Chuan Peng, Suhas Lohit, Joanna Matthiesen, Kevin Smith, and Joshua B Tenenbaum. 2024. Evaluating large vision-and-language models on children’s mathematical olympiads. *arXiv preprint arXiv:2406.15736*.

Ethan Chern, Haoyang Zou, Xuefeng Li, Jiewen Hu, Kehua Feng, Junlong Li, and Pengfei Liu. 2023. Generative ai for math: Abel. <https://github.com/GAIR-NLP/abel>.

Konstantin Chernyshev, Vitaliy Polshkov, Ekaterina Artemova, Alex Myasnikov, Vlad Stepanov, Alexei Miasnikov, and Sergei Tilga. 2024. U-math: A university-level benchmark for evaluating mathematical skills in llms. *arXiv preprint arXiv:2412.03205*.

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

Adrian Cosma, Ana-Maria Bucur, and Emilian Radoi. 2024. Romath: A mathematical reasoning benchmark in romanian. *arXiv preprint arXiv:2409.11074*.

777	Jiaxi Cui, Zongjian Li, Yang Yan, Bohua Chen, and	models using odyssey math data. arXiv preprint	832
778	Li Yuan. 2023. Chatlaw: Open-source legal large	arXiv:2406.18321 .	833
779	language model with integrated external knowledge		
780	bases. arXiv preprint arXiv:2306.16092 .		
781	Ishita Dasgupta, Christine Kaeser-Chen, Kenneth	Shengyu Feng, Xiang Kong, Shuang Ma, Aonan Zhang,	834
782	Marino, Arun Ahuja, Sheila Babayan, Felix Hill,	Dong Yin, Chong Wang, Ruoming Pang, and Yim-	835
783	and Rob Fergus. 2023. Collaborating with language	ing Yang. 2024. Step-by-step reasoning for math	836
784	models for embodied reasoning. arXiv preprint	problems via twisted sequential monte carlo. arXiv	837
785	arXiv:2302.00763 .	preprint arXiv:2410.01920 .	838
786	Arash Gholami Davoodi, Seyed Pouyan Mousavi	Deqing Fu, Ruohao Guo, Ghazal Khalighine-	839
787	Davoudi, and Pouya Pezeshkpour. 2024. Llms are	jad, Ollie Liu, Bhuwan Dhingra, Dani Yo-	840
788	not intelligent thinkers: Introducing mathematical	gatama, Robin Jia, and Willie Neiswanger. 2024a.	841
789	topic tree benchmark for comprehensive evaluation	Isobench: Benchmarking multimodal foundation	842
790	of llms. arXiv preprint arXiv:2406.05194 .	models on isomorphic representations. arXiv	843
791	Shumin Deng, Ningyu Zhang, Nay Oo, and Bryan	preprint arXiv:2404.01266 .	844
792	Hooi. 2023. Towards a unified view of answer cal-	Jiayi Fu, Lei Lin, Xiaoyang Gao, Pengli Liu, Zhengzong	845
793	ibration for multi-step reasoning. arXiv preprint	Chen, Zhirui Yang, Shengnan Zhang, Xue Zheng,	846
794	arXiv:2311.09101 .	Yan Li, Yuliang Liu, et al. 2023. Kwaiyimath: Tech-	847
795	Yihe Deng, Pan Lu, Fan Yin, Ziniu Hu, Sheng Shen,	nical report. arXiv preprint arXiv:2310.07488 .	848
796	James Zou, Kai-Wei Chang, and Wei Wang. 2024.	Xingyu Fu, Yushi Hu, Bangzheng Li, Yu Feng, Haoyu	849
797	Enhancing large vision language models with self-	Wang, Xudong Lin, Dan Roth, Noah A Smith, Wei-	850
798	training on image comprehension. arXiv preprint	Chiu Ma, and Ranjay Krishna. 2024b. Blink: Multi-	851
799	arXiv:2405.19716 .	modal large language models can see but not perceive.	852
800	Aniket Didolkar, Anirudh Goyal, Nan Rosemary Ke,	In European Conference on Computer Vision, pages	853
801	Siyuan Guo, Michal Valko, Timothy Lillicrap, Danilo	148–166. Springer.	854
802	Rezende, Yoshua Bengio, Michael Mozer, and San-	Bofei Gao, Feifan Song, Zhe Yang, Zefan Cai, Yibo	855
803	jeev Arora. 2024. Metacognitive capabilities of llms:	Miao, Qingxiu Dong, Lei Li, Chenghao Ma, Liang	856
804	An exploration in mathematical problem solving.	Chen, Runxin Xu, et al. 2024. Omni-math: A univer-	857
805	arXiv preprint arXiv:2405.12205 .	saral olympiad level mathematic benchmark for large	858
806	Prakhar Dixit and Tim Oates. 2024. Sbi-rag: Enhanc-	language models. arXiv preprint arXiv:2410.07985 .	859
807	ing math word problem solving for students through	Jiahui Gao, Renjie Pi, Jipeng Zhang, Jiacheng Ye, Wan-	860
808	schema-based instruction and retrieval-augmented	jun Zhong, Yufei Wang, Lanqing Hong, Jianhua Han,	861
809	generation. arXiv preprint arXiv:2410.13293 .	Hang Xu, Zhenguo Li, et al. 2023. G-llava: Solving	862
810	Duolingo. 2024. Duolingo official platfrom .	geometric problem with multi-modal large language	863
811	Zane Durante, Qiuyuan Huang, Naoki Wake, Ran Gong,	model. arXiv preprint arXiv:2312.11370 .	864
812	Jae Sung Park, Bidipta Sarkar, Rohan Taori, Yusuke	Philippe Gervais, Asya Fadeeva, and Andrii Maksai.	865
813	Noda, Demetri Terzopoulos, Yejin Choi, et al. 2024.	2024. Mathwriting: A dataset for handwritten	866
814	Agent ai: Surveying the horizons of multimodal in-	mathematical expression recognition. arXiv preprint	867
815	teraction. arXiv preprint arXiv:2401.03568 .	arXiv:2404.10690 .	868
816	Jingxuan Fan, Sarah Martinson, Erik Y Wang, Kaylie	Zhibin Gou, Zhihong Shao, Yeyun Gong, Yelong Shen,	869
817	Hausknecht, Jonah Brenner, Danxian Liu, Nianli	Yuju Yang, Minlie Huang, Nan Duan, and Weizhu	870
818	Peng, Corey Wang, and Michael P Brenner. 2024a.	Chen. 2023. Tora: A tool-integrated reasoning agent	871
819	Hardmath: A benchmark dataset for challenging	for mathematical problem solving. arXiv preprint	872
820	problems in applied mathematics. arXiv preprint	arXiv:2309.17452 .	873
821	arXiv:2410.09988 .	Aryan Gulati, Brando Miranda, Eric Chen, Emily	874
822	Wenqi Fan, Yujuan Ding, Liangbo Ning, Shijie Wang,	Xia, Kai Fronsdal, Bruno de Moraes Dumont, and	875
823	Hengyun Li, Dawei Yin, Tat-Seng Chua, and Qing	Sanmi Koyejo. 2024. Putnam-axiom: A functional	876
824	Pi. 2024b. A survey on rag meeting llms: To-	and static benchmark for measuring higher level	877
825	wards retrieval-augmented large language models. In	mathematical reasoning. In <i>The 4th Workshop on</i>	878
826	<i>Proceedings of the 30th ACM SIGKDD Conference</i>	<i>Mathematical Reasoning and AI at NeurIPS'24</i> .	879
827	<i>on Knowledge Discovery and Data Mining</i> , pages	Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song,	880
828	6491–6501.	Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma,	881
829	Meng Fang, Xiangpeng Wan, Fei Lu, Fei Xing, and	Peiyi Wang, Xiao Bi, et al. 2025. Deepseek-r1: In-	882
830	Kai Zou. 2024. Mathodyssey: Benchmarking math-	centivizing reasoning capability in llms via reinforce-	883
831	ematical problem-solving skills in large language	ment learning. arXiv preprint arXiv:2501.12948 .	884

885	Siyan Guo, Aniket Didolkar, Nan Rosemary Ke, Anirudh Goyal, Ferenc Huszár, and Bernhard Schölkopf. 2024a. Learning beyond pattern matching? assaying mathematical understanding in llms. arXiv preprint arXiv:2405.15485 .	942
886		943
887		944
888		945
889		
890	Taicheng Guo, Xiuying Chen, Yaqi Wang, Ruidi Chang, Shichao Pei, Nitesh V Chawla, Olaf Wiest, and Xiangliang Zhang. 2024b. Large language model based multi-agents: A survey of progress and challenges. arXiv preprint arXiv:2402.01680 .	946
891		947
892		948
893		949
894		
895	Tiezhen Guo, Qingwen Yang, Chen Wang, Yanyi Liu, Pan Li, Jiawei Tang, Dapeng Li, and Yingyou Wen. 2024c. Knowledgenavigator: Leveraging large language models for enhanced reasoning over knowledge graph. <i>Complex & Intelligent Systems</i> , 10(5):7063–7076.	950
896		951
897		952
898		953
899		954
900		
901	Himanshu Gupta, Shreyas Verma, Ujjwala Ananthaswaran, Kevin Scaria, Mihir Parmar, Swaroop Mishra, and Chitta Baral. 2024. Polymath: A challenging multi-modal mathematical reasoning benchmark. arXiv preprint arXiv:2410.14702 .	955
902		956
903		957
904		958
905		959
906	Tanmay Gupta and Aniruddha Kembhavi. 2023. Visual programming: Compositional visual reasoning without training. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 14953–14962.	960
907		961
908		962
909		963
910		964
911	Vernon Toh Yan Han, Ratish Puduppully, and Nancy F Chen. 2023. Veritymath: Advancing mathematical reasoning by self-verification through unit consistency. arXiv preprint arXiv:2311.07172 .	965
912		966
913		967
914		968
915		969
916	Xiaotian Han, Yiren Jian, Xuefeng Hu, Haogeng Liu, Yiqi Wang, Qihang Fan, Yuang Ai, Huaibo Huang, Ran He, Zhenheng Yang, et al. 2024. Infimwebmath-40b: Advancing multimodal pre-training for enhanced mathematical reasoning. arXiv preprint arXiv:2409.12568 .	970
917		971
918		972
919		973
920		974
921	Chaoqun He, Renjie Luo, Yuzhuo Bai, Shengding Hu, Zhen Leng Thai, Junhao Shen, Jinyi Hu, Xu Han, Yujie Huang, Yuxiang Zhang, et al. 2024. Olympiadbench: A challenging benchmark for promoting agi with olympiad-level bilingual multimodal scientific problems. arXiv preprint arXiv:2402.14008 .	975
922		976
923		977
924		978
925		979
926		980
927	Stephen J Hegedus and David O Tall. 2015. Foundations for the future: The potential of multimodal technologies for learning mathematics. In <i>Handbook of international research in mathematics education</i> , pages 543–562. Routledge.	981
928		982
929		983
930		984
931		
932	Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021. Measuring mathematical problem solving with the math dataset. arXiv preprint arXiv:2103.03874 .	985
933		986
934		987
935		988
936		989
937	Yushi Hu, Weijia Shi, Xingyu Fu, Dan Roth, Mari Ostendorf, Luke Zettlemoyer, Noah A Smith, and Ranjay Krishna. 2024. Visual sketchpad: Sketching as a visual chain of thought for multimodal language models. arXiv preprint arXiv:2406.09403 .	990
938		991
939		992
940		993
941		994
	Litian Huang, Xinguo Yu, Feng Xiong, Bin He, Shengbing Tang, and Jiawen Fu. 2024a. Hologram reasoning for solving algebra problems with geometry diagrams. arXiv preprint arXiv:2408.10592 .	
	Xuhan Huang, Qingning Shen, Yan Hu, Anningzhe Gao, and Benyou Wang. 2024b. Mamo: a mathematical modeling benchmark with solvers. arXiv preprint arXiv:2405.13144 .	
	Yiming Huang, Xiao Liu, Yeyun Gong, Zhibin Gou, Yelong Shen, Nan Duan, and Weizhu Chen. 2024c. Key-point-driven data synthesis with its enhancement on mathematical reasoning. arXiv preprint arXiv:2403.02333 .	
	Zihan Huang, Tao Wu, Wang Lin, Shengyu Zhang, Jingyuan Chen, and Fei Wu. 2024d. Autogeo: Automating geometric image dataset creation for enhanced geometry understanding. arXiv preprint arXiv:2409.09039 .	
	Nourhan Ibrahim, Samar Aboulela, Ahmed Ibrahim, and Rasha Kashef. 2024. A survey on augmenting knowledge graphs (kgs) with large language models (llms): models, evaluation metrics, benchmarks, and challenges. <i>Discover Artificial Intelligence</i> , 4(1):76.	
	Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. 2024. Openai o1 system card. arXiv preprint arXiv:2412.16720 .	
	Mengzhao Jia, Zhihan Zhang, Wenhao Yu, Fangkai Jiao, and Meng Jiang. 2024. Describe-then-reason: Improving multimodal mathematical reasoning through visual comprehension training. arXiv preprint arXiv:2404.14604 .	
	Hyoungwook Jin, Yoonsu Kim, Yeon Su Park, Bekzat Tilekbay, Jinho Son, and Juho Kim. 2024. Using large language models to diagnose math problem-solving skills at scale. In <i>Proceedings of the Eleventh ACM Conference on Learning@ Scale</i> , pages 471–475.	
	Bert Jonsson, Julia Mossegård, Johan Lithner, and Linnea Karlsson Wirebring. 2022. Creative mathematical reasoning: Does need for cognition matter? <i>Frontiers in Psychology</i> , 12:797807.	
	Sofia Jusslin, Kaisa Korpinen, Niina Lilja, Rose Martin, Johanna Lehtinen-Schnabel, and Eeva Anttila. 2022. Embodied learning and teaching approaches in language education: A mixed studies review. <i>Educational Research Review</i> , 37:100480.	
	Jikun Kang, Xin Zhe Li, Xi Chen, Amirreza Kazemi, Qianyi Sun, Boxing Chen, Dong Li, Xu He, Quan He, Feng Wen, et al. 2024. Mindstar: Enhancing math reasoning in pre-trained llms at inference time. arXiv preprint arXiv:2405.16265 .	

1106	Zhenwen Liang, Ye Liu, Tong Niu, Xiangliang Zhang, Yingbo Zhou, and Semih Yavuz. 2024b. Improving llm reasoning through scaling inference computation with collaborative verification. arXiv preprint arXiv:2410.05318 .	Xiao Liu, Zirui Wu, Xueqing Wu, Pan Lu, Kai-Wei Chang, and Yansong Feng. 2024d. Are llms capable of data-based statistical and causal reasoning? benchmarking advanced quantitative reasoning with data. arXiv preprint arXiv:2402.17644 .	1160
1107			1161
1108			1162
1109			1163
1110			1164
1111	Zhenwen Liang, Tianyu Yang, Jipeng Zhang, and Xiangliang Zhang. 2023a. Unimath: A foundational and multimodal mathematical reasoner. In Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing , pages 7126–7133.	Pan Lu, Hritik Bansal, Tony Xia, Jiacheng Liu, Chunyuan Li, Hannaneh Hajishirzi, Hao Cheng, Kai-Wei Chang, Michel Galley, and Jianfeng Gao. 2023. Mathvista: Evaluating mathematical reasoning of foundation models in visual contexts. arXiv preprint arXiv:2310.02255 .	1165
1112			1166
1113			1167
1114			1168
1115			1169
1116	Zhenwen Liang, Dian Yu, Xiaoman Pan, Wenlin Yao, Qingkai Zeng, Xiangliang Zhang, and Dong Yu. 2023b. Mint: Boosting generalization in mathematical reasoning via multi-view fine-tuning. arXiv preprint arXiv:2307.07951 .	Pan Lu, Ran Gong, Shibiao Jiang, Liang Qiu, Siyuan Huang, Xiaodan Liang, and Song-Chun Zhu. 2021. Inter-gps: Interpretable geometry problem solving with formal language and symbolic reasoning. arXiv preprint arXiv:2105.04165 .	1171
1117			1172
1118			1173
1119			1174
1120			1175
1121	Zhenwen Liang, Dian Yu, Wenhao Yu, Wenlin Yao, Zhihan Zhang, Xiangliang Zhang, and Dong Yu. 2024c. Mathchat: Benchmarking mathematical reasoning and instruction following in multi-turn interactions. arXiv preprint arXiv:2405.19444 .	Pan Lu, Baolin Peng, Hao Cheng, Michel Galley, Kai-Wei Chang, Ying Nian Wu, Song-Chun Zhu, and Jianfeng Gao. 2024a. Chameleon: Plug-and-play compositional reasoning with large language models. Advances in Neural Information Processing Systems , 36.	1176
1122			1177
1123			1178
1124			1179
1125			1180
1126	Zhenwen Liang, Jipeng Zhang, Lei Wang, Wei Qin, Yunshi Lan, Jie Shao, and Xiangliang Zhang. 2021. Mwp-bert: Numeracy-augmented pre-training for math word problem solving. arXiv preprint arXiv:2107.13435 .	Pan Lu, Liang Qiu, Kai-Wei Chang, Ying Nian Wu, Song-Chun Zhu, Tanmay Rajpurohit, Peter Clark, and Ashwin Kalyan. 2022a. Dynamic prompt learning via policy gradient for semi-structured mathematical reasoning. arXiv preprint arXiv:2209.14610 .	1182
1127			1183
1128			1184
1129			1185
1130			1186
1131	Zhenghao Lin, Zhibin Gou, Yeyun Gong, Xiao Liu, Yelong Shen, Ruochen Xu, Chen Lin, Yujia Yang, Jian Jiao, Nan Duan, et al. 2024. Rho-1: Not all tokens are what you need. arXiv preprint arXiv:2404.07965 .	Pan Lu, Liang Qiu, Wenhao Yu, Sean Welleck, and Kai-Wei Chang. 2022b. A survey of deep learning for mathematical reasoning. arXiv preprint arXiv:2212.10535 .	1187
1132			1188
1133			1189
1134			1190
1135	Haoxiong Liu, Yifan Zhang, Yifan Luo, and Andrew Chi-Chih Yao. 2024a. Augmenting math word problems via iterative question composing. arXiv preprint arXiv:2401.09003 .	Zimu Lu, Aojun Zhou, Houxing Ren, Ke Wang, Weikang Shi, Juntao Pan, Mingjie Zhan, and Hongsheng Li. 2024b. Mathgenie: Generating synthetic data with question back-translation for enhancing mathematical reasoning of llms. arXiv preprint arXiv:2402.16352 .	1191
1136			1192
1137			1193
1138			1194
1139	Hongwei Liu, Zilong Zheng, Yuxuan Qiao, Haodong Duan, Zhiwei Fei, Fengzhe Zhou, Wenwei Zhang, Songyang Zhang, Dahua Lin, and Kai Chen. 2024b. Mathbench: Evaluating the theory and application proficiency of llms with a hierarchical mathematics benchmark. arXiv preprint arXiv:2405.12209 .	Zimu Lu, Aojun Zhou, Ke Wang, Houxing Ren, Weikang Shi, Juntao Pan, Mingjie Zhan, and Hongsheng Li. 2024c. Mathcoder2: Better math reasoning from continued pretraining on model-translated mathematical code . Preprint, arXiv:2410.08196.	1195
1140			1196
1141			1197
1142			1198
1143			1199
1144			1200
1145	Junling Liu, Ziming Wang, Qichen Ye, Dading Chong, Peilin Zhou, and Yining Hua. 2023a. Qilin-med-vl: Towards chinese large vision-language model for general healthcare. arXiv preprint arXiv:2310.17956 .	Haipeng Luo, Qingfeng Sun, Can Xu, Pu Zhao, Jianguang Lou, Chongyang Tao, Xiubo Geng, Qingwei Lin, Shifeng Chen, and Dongmei Zhang. 2023. Wizardmath: Empowering mathematical reasoning for large language models via reinforced evol-instruct. arXiv preprint arXiv:2308.09583 .	1202
1146			1203
1147			1204
1148			1205
1149	Wentao Liu, Hanglei Hu, Jie Zhou, Yuyang Ding, Junsong Li, Jiayi Zeng, Mengliang He, Qin Chen, Bo Jiang, Aimin Zhou, et al. 2023b. Mathematical language models: A survey. arXiv preprint arXiv:2312.07622 .	Jingkun Ma, Runzhe Zhan, Derek F Wong, Yang Li, Di Sun, Hou Pong Chan, and Lidia S Chao. 2024. Vidsaidmath: Benchmarking visual-aided mathematical reasoning. arXiv preprint arXiv:2410.22995 .	1206
1150			1207
1151			1208
1152			1209
1153			1210
1154	Wentao Liu, Qianjun Pan, Yi Zhang, Zhuo Liu, Ji Wu, Jie Zhou, Aimin Zhou, Qin Chen, Bo Jiang, and Liang He. 2024c. Cmm-math: A chinese multimodal math dataset to evaluate and enhance the mathematics reasoning of large multimodal models. arXiv preprint arXiv:2409.02834 .	Yujun Mao, Yoon Kim, and Yilun Zhou. 2024. Champ: A competition-level dataset for fine-grained analyses of llms’ mathematical reasoning capabilities. arXiv preprint arXiv:2401.06961 .	1211
1155			1212
1156			1213
1157			1214
1158			1215
1159			

1216	Iman Mirzadeh, Keivan Alizadeh, Hooman Shahrokhi,	Jinghui Qin, Zhicheng Yang, Jiaqi Chen, Xiaodan Liang,	1267
1217	Oncel Tuzel, Samy Bengio, and Mehrdad Farajtabar.	and Liang Lin. 2023. Template-based contrastive	1268
1218	2024. Gsm-symbolic: Understanding the limitations	distillation pretraining for math word problem solv-	1269
1219	of mathematical reasoning in large language models.	ing. <i>IEEE Transactions on Neural Networks and</i>	1270
1220	arXiv preprint arXiv:2410.05229 .	<i>Learning Systems</i> .	1271
1221	Swaroop Mishra, Matthew Finlayson, Pan Lu, Leonard	Qwen. 2024. Qwen2-math technical report .	1272
1222	Tang, Sean Welleck, Chitta Baral, Tanmay Rajpuro-	AM Rahman, Junyi Ye, Wei Yao, Wenpeng Yin, and	1273
1223	hit, Oyvind Tafjord, Ashish Sabharwal, Peter Clark,	Guiling Wang. 2024. From blind solvers to logi-	1274
1224	et al. 2022. Lila: A unified benchmark for mathemat-	cal thinkers: Benchmarking llms’ logical integrity	1275
1225	ical reasoning. arXiv preprint arXiv:2210.17517 .	on faulty mathematical problems. arXiv preprint	1276
1226	MistralAI. 2024. Mathstral official platform .	arXiv:2410.18921 .	1277
1227	MoonshotAI. 2024. k0-math official platform .	Machel Reid, Nikolay Savinov, Denis Teplyashin,	1278
1228	Harold Mouchere, Christian Viard-Gaudin, Dae Hwan	Dmitry Lepikhin, Timothy Lillicrap, Jean-baptiste	1279
1229	Kim, Jin Hyung Kim, and Utpal Garain. 2011.	Alayrac, Radu Soricut, Angeliki Lazaridou, Orhan Fi-	1280
1230	Crohme2011: Competition on recognition of on-	rat, Julian Schrittwieser, et al. 2024. Gemini 1.5: Un-	1281
1231	line handwritten mathematical expressions. In <i>2011</i>	locking multimodal understanding across millions of	1282
1232	<i>international conference on document analysis and</i>	tokens of context. arXiv preprint arXiv:2403.05530 .	1283
1233	<i>recognition</i> , pages 1497–1500. IEEE.	Ankit Satpute, Noah Gießing, André Greiner-Petter,	1284
1234	Niklas Muennighoff, Zitong Yang, Weijia Shi, Xi-	Moritz Schubotz, Olaf Teschke, Akiko Aizawa, and	1285
1235	ang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke	Bela Gipp. 2024. Can llms master math? investigat-	1286
1236	Zettlemoyer, Percy Liang, Emmanuel Candès, and	ing large language models on math stack exchange.	1287
1237	Tatsunori Hashimoto. 2025. s1: Simple test-time	In <i>Proceedings of the 47th International ACM</i>	1288
1238	scaling. arXiv preprint arXiv:2501.19393 .	<i>SIGIR Conference on Research and Development in</i>	1289
1239	Zabir Al Nazi and Wei Peng. 2024. Large language	<i>Information Retrieval</i> , pages 2316–2320.	1290
1240	models in healthcare and medical domain: A review.	Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu,	1291
1241	In <i>Informatics</i> , volume 11, page 57. MDPI.	Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan	1292
1242	Bolin Ni, JingCheng Hu, Yixuan Wei, Houwen Peng,	Zhang, YK Li, Y Wu, et al. 2024. Deepseekmath:	1293
1243	Zheng Zhang, Gaofeng Meng, and Han Hu. 2024.	Pushing the limits of mathematical reasoning in open	1294
1244	Xwin-lm: Strong and scalable alignment practice for	language models. arXiv preprint arXiv:2402.03300 .	1295
1245	llms. arXiv preprint arXiv:2405.20335 .	Wenhao Shi, Zhiqiang Hu, Yi Bin, Junhua Liu, Yang	1296
1246	OpenAI. 2024. Gpt-4o system card .	Yang, See-Kiong Ng, Lidong Bing, and Roy Ka-Wei	1297
1247	Kay L O’Halloran. 2015. The language of learn-	Lee. 2024. Math-llava: Bootstrapping mathemati-	1298
1248	ing mathematics: A multimodal perspective. <i>The</i>	cal reasoning for multimodal large language models.	1299
1249	<i>Journal of Mathematical Behavior</i> , 40:63–74.	arXiv preprint arXiv:2406.17294 .	1300
1250	Jeff Z Pan, Simon Razniewski, Jan-Christoph Kalo,	Shiven Sinha, Ameya Prabhu, Ponnurangam Ku-	1301
1251	Sneha Singhania, Jiaoyan Chen, Stefan Dietze, Ha-	maraguru, Siddharth Bhat, and Matthias Bethge.	1302
1252	jira Jabeen, Janna Omeliiyanenko, Wen Zhang, Mat-	2024. Wu’s method can boost symbolic ai to ri-	1303
1253	teo Lissandrini, et al. 2023. Large language models	val silver medalists and alphageometry to outper-	1304
1254	and knowledge graphs: Opportunities and challenges.	form gold medalists at imo geometry. arXiv preprint	1305
1255	arXiv preprint arXiv:2308.06374 .	arXiv:2404.06405 .	1306
1256	Shuai Peng, Di Fu, Liangcai Gao, Xiuqin Zhong, Hong-	Shezheng Song, Xiaopeng Li, Shasha Li, Shan Zhao,	1307
1257	guang Fu, and Zhi Tang. 2024. Multimath: Bridging	Jie Yu, Jun Ma, Xiaoguang Mao, and Weimin Zhang.	1308
1258	visual and mathematical reasoning for large language	2023. How to bridge the gap between modalities: A	1309
1259	models. arXiv preprint arXiv:2409.00147 .	comprehensive survey on multimodal large language	1310
1260	Gabriel Poesia, David Broman, Nick Haber, and	model. arXiv preprint arXiv:2311.07594 .	1311
1261	Noah D Goodman. 2024. Learning formal math-	SquirrelAiLearning. 2024. Squirrel ai official platfrom .	1312
1262	ematics from intrinsic motivation. arXiv preprint	Pragya Srivastava, Manuj Malik, Vivek Gupta, Tanuja	1313
1263	arXiv:2407.00695 .	Ganu, and Dan Roth. 2024. Evaluating llms’ math-	1314
1264	Stanislas Polu and Ilya Sutskever. 2021. Generative	ematical reasoning in financial document question	1315
1265	language modeling for automated theorem proving.	answering. In <i>Findings of the Association for</i>	1316
1266	arXiv preprint arXiv:2009.03393 .	<i>Computational Linguistics ACL 2024</i> , pages 3853–	1317
		3878.	1318
		Kai Sun, Yushi Bai, Ji Qi, Lei Hou, and Juanzi Li.	1319
		2024a. Mm-math: Advancing multimodal math	1320
		evaluation with process evaluation and fine-grained	1321

1322	classification. In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 1358–1375.	1377
1323		1378
1324		1379
1325	Liangtai Sun, Yang Han, Zihan Zhao, Da Ma, Zhen-	1380
1326	nan Shen, Baocai Chen, Lu Chen, and Kai Yu.	
1327	2024b. Scieval: A multi-level large language model	1381
1328	evaluation benchmark for scientific research. In	1382
1329	<i>Proceedings of the AAAI Conference on Artificial</i>	1383
1330	<i>Intelligence</i> , volume 38, pages 19053–19061.	1384
1331		1385
1332	Lin Zhuang Sun, Hao Liang, Jingxuan Wei, Bihui Yu,	1386
1333	Conghui He, Zenan Zhou, and Wentao Zhang. 2024c.	
1334	Beats: Optimizing llm mathematical capabilities with	1387
1335	backverify and adaptive disambiguate based efficient	1388
	tree search. <i>arXiv preprint arXiv:2409.17972</i> .	1389
1336		1390
1337	Dídac Surís, Sachit Menon, and Carl Vondrick. 2023.	1391
1338	Vipergpt: Visual inference via python execution	
1339	for reasoning. In <i>Proceedings of the IEEE/CVF</i>	1392
1340	<i>International Conference on Computer Vision</i> , pages	1393
	11888–11898.	1394
1341	TALEducation. 2023. <i>Mathgpt official platform</i> .	1395
1342		1396
1343	Jiamin Tang, Chao Zhang, Xudong Zhu, and Mengchi	1397
1344	Liu. 2024a. Tangram: A challenging benchmark	1398
1345	for geometric element recognizing. <i>arXiv preprint</i>	
	<i>arXiv:2408.13854</i> .	1399
1346	Zhengyang Tang, Xingxing Zhang, Benyou Wang, and	1400
1347	Furu Wei. 2024b. Mathscale: Scaling instruction	1401
1348	tuning for mathematical reasoning. <i>arXiv preprint</i>	1402
1349	<i>arXiv:2403.02884</i> .	1403
1350		1404
1351	Amitayush Thakur, George Tsoukalas, Yeming Wen,	1405
1352	Jimmy Xin, and Swarat Chaudhuri. 2024. An in-	1406
1353	context learning agent for formal theorem-proving.	1407
	In <i>First Conference on Language Modeling</i> .	1408
1354		1409
1355	Yuxuan Tong, Xiwen Zhang, Rui Wang, Ruidong Wu,	
1356	and Junxian He. 2024. Dart-math: Difficulty-aware	1410
1357	rejection tuning for mathematical problem-solving.	1411
	<i>arXiv preprint arXiv:2407.13690</i> .	1412
1358	Shubham Toshniwal, Wei Du, Ivan Moshkov, Branislav	1413
1359	Kisacanian, Alexan Ayrapetyan, and Igor Gitman.	1414
1360	2024a. Openmathinstruct-2: Accelerating ai for math	1415
1361	with massive open-source instruction data. <i>arXiv</i>	1416
1362	<i>preprint arXiv:2410.01560</i> .	
1363	Shubham Toshniwal, Ivan Moshkov, Sean Narenthi-	1417
1364	ran, Daria Gitman, Fei Jia, and Igor Gitman. 2024b.	1418
1365	Openmathinstruct-1: A 1.8 million math instruction	1419
1366	tuning dataset. <i>arXiv preprint arXiv:2402.10176</i> .	1420
1367		1421
1368	Trieu H Trinh, Yuhuai Wu, Quoc V Le, He He,	
1369	and Thang Luong. 2024. Solving olympiad ge-	1422
1370	ometry without human demonstrations. <i>Nature</i> ,	1423
	625(7995):476–482.	1424
1371		1425
1372	George Tsoukalas, Jasper Lee, John Jennings, Jimmy	
1373	Xin, Michelle Ding, Michael Jennings, Amitayush	1426
1374	Thakur, and Swarat Chaudhuri. Putnambench: A	1427
1375	multilingual competition-mathematics benchmark for	1428
1376	formal theorem-proving. In <i>AI for Math Workshop@</i>	1429
	<i>ICML 2024</i> .	1430
	Ke Wang, Junting Pan, Weikang Shi, Zimu Lu, Mingjie	
	Zhan, and Hongsheng Li. 2024a. Measuring mul-	
	timodal mathematical reasoning with math-vision	
	dataset. <i>arXiv preprint arXiv:2402.14804</i> .	
	Ke Wang, Houxing Ren, Aojun Zhou, Zimu Lu,	
	Sichun Luo, Weikang Shi, Renrui Zhang, Linqi	
	Song, Mingjie Zhan, and Hongsheng Li. 2023a.	
	Mathcoder: Seamless code integration in llms for	
	enhanced mathematical reasoning. <i>arXiv preprint</i>	
	<i>arXiv:2310.03731</i> .	
	Lei Wang, Shan Dong, Yuhui Xu, Hanze Dong, Yalu	
	Wang, Amrita Saha, Ee-Peng Lim, Caiming Xiong,	
	and Doyen Sahoo. 2024b. Mathhay: An automated	
	benchmark for long-context mathematical reasoning	
	in llms. <i>arXiv preprint arXiv:2410.04698</i> .	
	Peiyi Wang, Lei Li, Zhihong Shao, Runxin Xu, Damai	
	Dai, Yifei Li, Deli Chen, Yu Wu, and Zhifang	
	Sui. 2024c. Math-shepherd: Verify and reinforce	
	llms step-by-step without human annotations. In	
	<i>Proceedings of the 62nd Annual Meeting of the</i>	
	<i>Association for Computational Linguistics (Volume</i>	
	<i>1: Long Papers)</i> , pages 9426–9439.	
	Shen Wang, Tianlong Xu, Hang Li, Chaoli Zhang,	
	Joleen Liang, Jiliang Tang, Philip S Yu, and Qing-	
	song Wen. 2024d. Large language models for ed-	
	ucation: A survey and outlook. <i>arXiv preprint</i>	
	<i>arXiv:2403.18105</i> .	
	Xiaoxuan Wang, Ziniu Hu, Pan Lu, Yanqiao Zhu,	
	Jieyu Zhang, Satyen Subramaniam, Arjun R Loomba,	
	Shichang Zhang, Yizhou Sun, and Wei Wang.	
	2023b. Scibench: Evaluating college-level scientific	
	problem-solving abilities of large language models.	
	<i>arXiv preprint arXiv:2307.10635</i> .	
	Yiqi Wang, Wentao Chen, Xiaotian Han, Xudong Lin,	
	Haiteng Zhao, Yongfei Liu, Bohan Zhai, Jianbo Yuan,	
	Quanzeng You, and Hongxia Yang. 2024e. Exploring	
	the reasoning abilities of multimodal large language	
	models (mllms): A comprehensive survey on emerg-	
	ing trends in multimodal reasoning. <i>arXiv preprint</i>	
	<i>arXiv:2401.06805</i> .	
	Yixu Wang, Wenpin Qian, Hong Zhou, Jianfeng Chen,	
	and Kai Tan. 2023c. Exploring new frontiers of deep	
	learning in legal practice: A case study of large lan-	
	guage models. <i>International Journal of Computer</i>	
	<i>Science and Information Technology</i> , 1(1):131–138.	
	Chenrui Wei, Mengzhou Sun, and Wei Wang.	
	2024. Proving olympiad algebraic inequalities	
	without human demonstrations. <i>arXiv preprint</i>	
	<i>arXiv:2406.14219</i> .	
	Jiayang Wu, Wensheng Gan, Zefeng Chen, Shicheng	
	Wan, and S Yu Philip. 2023a. Multimodal large lan-	
	guage models: A survey. In <i>2023 IEEE International</i>	
	<i>Conference on Big Data (BigData)</i> , pages 2247–	
	2256. IEEE.	

1431	Shijie Wu, Ozan Irsoy, Steven Lu, Vadim Dabravolski,	Hanyi Xu, Wensheng Gan, Zhenlian Qi, Jiayang	1484
1432	Mark Dredze, Sebastian Gehrmann, Prabhajan Kam-	Wu, and Philip S Yu. 2024a. Large language	1485
1433	badur, David Rosenberg, and Gideon Mann. 2023b.	models for education: A survey. arXiv preprint	1486
1434	Bloomberggpt: A large language model for finance.	arXiv:2405.13001 .	1487
1435	arXiv preprint arXiv:2303.17564 .		
1436	Ting Wu, Xuefeng Li, and Pengfei Liu. 2024a. Progress	Liang Xu, Hang Xue, Lei Zhu, and Kangkang Zhao.	1488
1437	or regress? self-improvement reversal in post-	2024b. Superclue-math6: Graded multi-step math	1489
1438	training. arXiv preprint arXiv:2407.05013 .	reasoning benchmark for llms in chinese. arXiv	1490
		preprint arXiv:2401.11819 .	1491
1439	Zhenyu Wu, Meng Jiang, and Chao Shen. 2024b. Get	Tianlong Xu, Richard Tong, Jing Liang, Xing Fan,	1492
1440	an a in math: Progressive rectification prompting. In	Haoyang Li, and Qingsong Wen. 2024c. Founda-	1493
1441	Proceedings of the AAAI Conference on Artificial	tion models for education: Promises and prospects.	1494
1442	Intelligence , volume 38, pages 19288–19296.	arXiv preprint arXiv:2405.10959 .	1495
1443	Zijian Wu, Suozhi Huang, Zhejian Zhou, Huaiyuan	Tianlong Xu, Yi-Fan Zhang, Zhendong Chu, Shen	1496
1444	Ying, Jiayu Wang, Dahua Lin, and Kai Chen. 2024c.	Wang, and Qingsong Wen. 2024d. Ai-driven virtual	1497
1445	Internlm2. 5-stepprover: Advancing automated theo-	teacher for enhanced educational efficiency: Leverag-	1498
1446	rem proving via expert iteration on large-scale lean	ing large pretrain models for autonomous error analy-	1499
1447	problems. arXiv preprint arXiv:2410.15700 .	sis and correction. arXiv preprint arXiv:2409.09403 .	1500
1448	Zhiheng Xi, Wenxiang Chen, Xin Guo, Wei He, Yiwen	Xin Xu, Tong Xiao, Zitong Chao, Zhenya Huang, Can	1501
1449	Ding, Boyang Hong, Ming Zhang, Junzhe Wang,	Yang, and Yang Wang. 2024e. Can llms solve	1502
1450	Senjie Jin, Enyu Zhou, et al. 2023. The rise and	longer math word problems better? arXiv preprint	1503
1451	potential of large language model based agents: A	arXiv:2405.14804 .	1504
1452	survey. arXiv preprint arXiv:2309.07864 .		
1453	Changrong Xiao, Sean Xin Xu, and Kunpeng Zhang.	Xin Xu, Jiabin Zhang, Tianhao Chen, Zitong Chao,	1505
1454	2023. Multimodal data augmentation for image cap-	Jishan Hu, and Can Yang. 2025b. Ugmathbench: A	1506
1455	tioning using diffusion models. In Proceedings of	diverse and dynamic benchmark for undergraduate-	1507
1456	the 1st Workshop on Large Generative Models Meet	level mathematical reasoning with large language	1508
1457	Multimodal Applications , pages 23–33.	models. arXiv preprint arXiv:2501.13766 .	1509
1458	Junlin Xie, Zhihong Chen, Ruifei Zhang, Xiang Wan,	Yifan Xu, Xiao Liu, Xinghan Liu, Zhenyu Hou, Yueyan	1510
1459	and Guanbin Li. 2024. Large multimodal agents: A	Li, Xiaohan Zhang, Zihan Wang, Aohan Zeng,	1511
1460	survey. arXiv preprint arXiv:2402.15116 .	Zhengxiao Du, Wenyi Zhao, et al. 2024f. Chatglm-	1512
		math: Improving math problem-solving in large lan-	1513
1461	Huajian Xin, Daya Guo, Zhihong Shao, Zhizhou Ren,	guage models with a self-critique pipeline. arXiv	1514
1462	Qihao Zhu, Bo Liu, Chong Ruan, Wenda Li, and	preprint arXiv:2404.02893 .	1515
1463	Xiaodan Liang. 2024a. Deepseek-prover: Advancing	Yibo Yan and Joey Lee. 2024. Georeasoner: Reason-	1516
1464	theorem proving in llms through large-scale synthetic	ing on geospatially grounded context for natural lan-	1517
1465	data. arXiv preprint arXiv:2405.14333 .	guage understanding. In Proceedings of the 33rd	1518
		ACM International Conference on Information and	1519
1466	Huajian Xin, ZZ Ren, Junxiao Song, Zhihong Shao,	Knowledge Management , pages 4163–4167.	1520
1467	Wanjia Zhao, Haocheng Wang, Bo Liu, Liyue Zhang,	Yibo Yan, Shen Wang, Jiahao Huo, Hang Li, Boyan Li,	1521
1468	Xuan Lu, Qiusi Du, et al. 2024b. Deepseek-prover-	Jiamin Su, Xiong Gao, Yi-Fan Zhang, Tianlong Xu,	1522
1469	v1. 5: Harnessing proof assistant feedback for re-	Zhendong Chu, et al. 2024a. Errorradar: Benchmark-	1523
1470	inforcement learning and monte-carlo tree search.	ing complex mathematical reasoning of multimodal	1524
1471	arXiv preprint arXiv:2408.08152 .	large language models via error detection. arXiv	1525
		preprint arXiv:2410.04509 .	1526
1472	Wei Xiong, Chengshuai Shi, Jiaming Shen, Aviv Rosen-	Yibo Yan, Shen Wang, Jiahao Huo, Jingheng Ye, Zhen-	1527
1473	berg, Zhen Qin, Daniele Calandriello, Misha Khal-	dong Chu, Xuming Hu, Philip S Yu, Carla Gomes,	1528
1474	man, Rishabh Joshi, Bilal Piot, Mohammad Saleh,	Bart Selman, and Qingsong Wen. 2025. Posi-	1529
1475	et al. 2024. Building math agents with multi-	tion: Multimodal large language models can signif-	1530
1476	turn iterative preference learning. arXiv preprint	icantly advance scientific reasoning. arXiv preprint	1531
1477	arXiv:2409.02392 .	arXiv:2502.02871 .	1532
1478	Fengli Xu, Qian Yue Hao, Zefang Zong, Jingwei Wang,	Yibo Yan, Haomin Wen, Siru Zhong, Wei Chen,	1533
1479	Yunke Zhang, Jingyi Wang, Xiaochong Lan, Jiahui	Haodong Chen, Qingsong Wen, Roger Zimmer-	1534
1480	Gong, Tianjian Ouyang, Fanjin Meng, et al. 2025a.	mann, and Yuxuan Liang. 2024b. Urbanclip: Learn-	1535
1481	Towards large reasoning models: A survey of rein-	ing text-enhanced urban region profiling with con-	1536
1482	forced reasoning with large language models. arXiv	trastive language-image pretraining from the web. In	1537
1483	preprint arXiv:2501.09686 .	Proceedings of the ACM on Web Conference 2024 ,	1538
		pages 4006–4017.	1539

1540	Yuchen Yan, Jin Jiang, Yang Liu, Yixin Cao, Xin Xu, Xunliang Cai, Jian Shao, et al. 2024c. S3c-math: Spontaneous step-level self-correction makes large language models better mathematical reasoners. arXiv preprint arXiv:2409.01524 .	1595
1541		1596
1542		1597
1543		1598
1544		1599
1545	An Yang, Beichen Zhang, Binyuan Hui, Bofei Gao, Bowen Yu, Chengpeng Li, Dayiheng Liu, Jianhong Tu, Jingren Zhou, Junyang Lin, et al. 2024a. Qwen2.5-math technical report: Toward mathematical expert model via self-improvement. arXiv preprint arXiv:2409.12122 .	1600
1546		1601
1547		1602
1548		1603
1549		
1550		
1551	Hongyang Yang, Xiao-Yang Liu, and Christina Dan Wang. 2023a. Fingpt: Open-source financial large language models. arXiv preprint arXiv:2306.06031 .	1604
1552		1605
1553		1606
1554		1607
1555		1608
1556	Zhen Yang, Jinhao Chen, Zhengxiao Du, Wenmeng Yu, Weihao Wang, Wenyi Hong, Zhihuan Jiang, Bin Xu, Yuxiao Dong, and Jie Tang. 2024b. Mathglm-vision: Solving mathematical problems with multimodal large language model. arXiv preprint arXiv:2409.13729 .	1609
1557		1610
1558		1611
1559		1612
1560		1613
1561	Zhen Yang, Ming Ding, Qingsong Lv, Zhihuan Jiang, Zehai He, Yuyi Guo, Jinfeng Bai, and Jie Tang. 2023b. Gpt can solve mathematical problems without a calculator. arXiv preprint arXiv:2309.03241 .	1614
1562		1615
1563		1616
1564		
1565	Zhicheng Yang, Jinghui Qin, Jiaqi Chen, Liang Lin, and Xiaodan Liang. 2022. Logicsolver: Towards interpretable math word problem solving with logical prompt-enhanced learning. arXiv preprint arXiv:2205.08232 .	1617
1566		1618
1567		1619
1568		1620
1569		1621
1570	Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. 2022. React: Synergizing reasoning and acting in language models. arXiv preprint arXiv:2210.03629 .	1622
1571		
1572		
1573		
1574	Shukang Yin, Chaoyou Fu, Sirui Zhao, Ke Li, Xing Sun, Tong Xu, and Enhong Chen. 2024. A survey on multimodal large language models. National Science Review , page nwae403.	1623
1575		1624
1576		1625
1577		1626
1578		1627
1579	Huaiyuan Ying, Shuo Zhang, Linyang Li, Zhejian Zhou, Yunfan Shao, Zhaoye Fei, Yichuan Ma, Jiawei Hong, Kuikun Liu, Ziyi Wang, et al. 2024. Internlm-math: Open math large language models toward verifiable reasoning. arXiv preprint arXiv:2402.06332 .	1628
1580		1629
1581		1630
1582		1631
1583	Dian Yu, Baolin Peng, Ye Tian, Linfeng Song, Haitao Mi, and Dong Yu. 2024a. Siam: Self-improving code-assisted mathematical reasoning of large language models. arXiv preprint arXiv:2408.15565 .	1632
1584		1633
1585		1634
1586		1635
1587	Fei Yu, Hongbo Zhang, Prayag Tiwari, and Benyou Wang. 2024b. Natural language reasoning, a survey. ACM Computing Surveys , 56(12):1–39.	1636
1588		1637
1589		
1590	Longhui Yu, Weisen Jiang, Han Shi, Jincheng Yu, Zhengying Liu, Yu Zhang, James T Kwok, Zhengguo Li, Adrian Weller, and Weiyang Liu. 2023. Metamath: Bootstrap your own mathematical questions for large language models. arXiv preprint arXiv:2309.12284 .	1638
1591		1639
1592		1640
1593		1641
1594		1642
	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 Wenhao Chen. 2024a. Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi. In Proceedings of CVPR .	1643
		1644
		1645
		1646
	Xiang Yue, Xingwei Qu, Ge Zhang, Yao Fu, Wenhao Huang, Huan Sun, Yu Su, and Wenhao Chen. 2023. Mammoth: Building math generalist models through hybrid instruction tuning. arXiv preprint arXiv:2309.05653 .	1647
		1648
		1649
		1650
	Xiang Yue, Tianyu Zheng, Yuansheng Ni, Yubo Wang, Kai Zhang, Shengbang Tong, Yuxuan Sun, Botao Yu, Ge Zhang, Huan Sun, et al. 2024b. Mmmu-pro: A more robust multi-discipline multimodal understanding benchmark. arXiv preprint arXiv:2409.02813 .	1647
		1648
		1649
		1650
	Xiang Yue, Tuney Zheng, Ge Zhang, and Wenhao Chen. 2024c. Mammoth2: Scaling instructions from the web. arXiv preprint arXiv:2405.03548 .	1647
		1648
		1649
		1650
	Liang Zeng, Liangjun Zhong, Liang Zhao, Tianwen Wei, Liu Yang, Jujie He, Cheng Cheng, Rui Hu, Yang Liu, Shuicheng Yan, et al. 2024. Skywork-math: Data scaling laws for mathematical reasoning in large language models—the story goes on. arXiv preprint arXiv:2407.08348 .	1647
		1648
		1649
		1650
	Beichen Zhang, Kun Zhou, Xilin Wei, Xin Zhao, Jing Sha, Shijin Wang, and Ji-Rong Wen. 2024a. Evaluating and improving tool-augmented computation-intensive math reasoning. Advances in Neural Information Processing Systems , 36.	1647
		1648
		1649
		1650
	Di Zhang, Jianbo Wu, Jingdi Lei, Tong Che, Jia-tong Li, Tong Xie, Xiaoshui Huang, Shufei Zhang, Marco Pavone, Yuqiang Li, et al. 2024b. Llamaberry: Pairwise optimization for o1-like olympiad-level mathematical reasoning. arXiv preprint arXiv:2410.02884 .	1647
		1648
		1649
		1650
	Duzhen Zhang, Yahan Yu, Jiahua Dong, Chenxing Li, Dan Su, Chenhui Chu, and Dong Yu. 2024c. Mm-llms: Recent advances in multimodal large language models. arXiv preprint arXiv:2401.13601 .	1647
		1648
		1649
		1650
	Jiaxin Zhang, Zhongzhi Li, Mingliang Zhang, Fei Yin, Chenglin Liu, and Yashar Moshfeghi. 2024d. Geoeval: benchmark for evaluating llms and multi-modal models on geometry problem-solving. arXiv preprint arXiv:2402.10104 .	1647
		1648
		1649
		1650
	Mengxue Zhang, Zichao Wang, Zhichao Yang, Weiqi Feng, and Andrew Lan. 2023. Interpretable math word problem solution generation via step-by-step planning. arXiv preprint arXiv:2306.00784 .	1647
		1648
		1649
		1650
	Ming-Liang Zhang, Zhong-Zhi Li, Fei Yin, Liang Lin, and Cheng-Lin Liu. 2024e. Fuse, reason and verify: Geometry problem solving with parsed clauses from diagram. arXiv preprint arXiv:2407.07327 .	1647
		1648
		1649
		1650

1651	Renrui Zhang, Xinyu Wei, Dongzhi Jiang, Ziyu Guo, Shicheng Li, Yichi Zhang, Chengzhuo Tong, Jiaming Liu, Aojun Zhou, Bin Wei, et al. 2024g. Mavis: Mathematical visual instruction tuning with an automatic data engine. arXiv preprint arXiv:2407.08739 .	1707
1652		1708
1653		1709
1654		1710
1655		1711
1656		1712
1657		1713
1658		1714
1659		1715
1660		1716
1661		1717
1662		1718
1663		1719
1664		1720
1665		1721
1666		1722
1667		1723
1668		1724
1669		1725
1670		1726
1671		1727
1672		1728
1673		1729
1674		1730
1675		1731
1676		1732
1677		1733
1678		1734
1679		1735
1680		1736
1681		1737
1682		1738
1683		1739
1684		1740
1685		1741
1686		1742
1687		1743
1688		1744
1689		1745
1690		1746
1691		1747
1692		1748
1693		1749
1694		1750
1695		1751
1696		1752
1697		1753
1698		1754
1699		1755
1700		
1701		
1702		
1703		
1704		
1705		
1706		

A Details of Math-LLMs' Progress

The rapid development of general-purpose LLMs has made significant advancements in natural language processing tasks. However, the development of domain-specific models remains a core requirement, as they are better equipped to handle specialized tasks that general models may not address effectively. This is particularly true in fields such as healthcare (Liu et al., 2023a; Nazi and Peng, 2024), law (Cui et al., 2023; Zhou et al., 2024c; Wang et al., 2023c), finance (Wu et al., 2023b; Yang et al., 2023a; Zhang and Yang, 2023), and urban science (Yan et al., 2024b; Zou et al., 2025; Yan and Lee, 2024), where domain-specific knowledge is critical for high accuracy and performance.

In the case of mathematical reasoning, general models may struggle with tasks that require a deep understanding of complex mathematical concepts, structures, and problem-solving steps (Yan et al., 2025). Therefore, the development of math-specific LLMs is of paramount importance, as these models are designed to enhance performance in mathematical reasoning, theorem proving, equation solving, and other math-intensive tasks.

Therefore, Table 4 provides a detailed overview of various math-specific LLMs (*i.e.*, Math-(LLMs)), sorted by their release date. It includes information about the organization behind each model, the release date, publication details, language(s) supported, parameter size, evaluation benchmarks, and whether the model is open source.

Key findings are summarized as follows:

- 1. Release Trends:** The models started emerging in 2020, with a significant increase in the number of releases from 2022 onward, indicating a growing interest in developing math-specific LLMs.
- 2. Parameter Sizes:** There is a noticeable trend towards larger parameter sizes, with some models offering up to 130B parameters, reflecting the increasing computational capacity for handling complex mathematical tasks.
- 3. Evaluation Benchmarks:** Many models are evaluated on popular benchmarks like GSM8K, MATH, and MMLU, highlighting the focus on improving performance across well-established mathematical reasoning datasets.
- 4. Multilingual Support:** While most models are focused on English, a few (*e.g.*, MathGPT & Math-LLM) also support Chinese, showing a trend towards multilingual capabilities.
- 5. Open Source:** A significant number of models are open-source, allowing broader access and fostering further research and development in the field.

In summary, the table reflects the rapid development of specialized Math-LLMs, with an increasing trend towards larger models, comprehensive evaluation benchmarks, and support for multilingual applications.

B Summary of Benchmarks

Table 3 summarizes the LLM-based benchmarks for mathematical reasoning.

C Illustration of More Cases

Figure 5 illustrates the diverse multimodal cases of mathematical reasoning settings.

C.1 Multimodal Plane Geometry Setting

The Multimodal Plane Geometry Setting involves mathematical problems that require understanding and reasoning about 2D geometric relationships. These problems typically focus on fundamental geometric concepts, such as points, lines, angles, and triangles, often leveraging trigonometric principles like sine, cosine, or tangent. Visually, these questions are characterized by clear plane diagrams with labeled points, angles, and lengths. Students need to interpret these visuals to solve for unknown distances, angles, or other parameters. The defining feature here is the emphasis on 2D spatial relationships and the need to derive solutions from diagrammatic representations that combine measurements and geometry.

C.2 Multimodal Solid Geometry Setting

The Multimodal Solid Geometry Setting shifts the focus from 2D to 3D shapes and figures, such as cylinders, spheres, cubes, or cones. These questions often require students to compute surface area, volume, or height based on given measurements or constraints. Visually, these questions feature 3D diagrams with dimensions like radius, height, or length, typically annotated on the figure to help guide problem-solving. The main distinction is

Benchmarks	Venue	Language	Size	Source	Level(s)	Evaluation	Model(s)	Task(s)
DynaMath (Zou et al., 2024) ★	ICLR'25	English	5,010	S P G	E M H U	Both	Closed/Open	S
MathCheck (Zhou et al., 2024d) ★	ICLR'25	English/Chinese	4,536	P	E M H U	Discriminative	Closed/Open/Math	S
GSM-Symbolic (Mirzadeh et al., 2024)	ICLR'25	English	5,000	P	E	Discriminative	Closed/Open	S
Omni-MATH (Gao et al., 2024)	ICLR'25	English	4,428	S	C	Discriminative	Closed/Open/Math	S
HARDMath (Fan et al., 2024a)	ICLR'25	English	1,466	G	U	Both	Closed/Open	S
OpenMathInstruct-2 (Toshniwal et al., 2024a)	ICLR'25	English	14,000,000	P G	E H C	Discriminative	Open/Math	S
UGMathBench (Xu et al., 2025b)	ICLR'25	English	5,062	S	U	Discriminative	Closed/Open/Math	S
M ³ CoT _{math} (Chen et al., 2024b) ★	ACL'24	English	1,166	P G	C	Discriminative	Closed/Open	S
GSM-Plus (Li et al., 2024d)	ACL'24	English	10,552	P	E M H U	Generative	Closed/Open/Math	S
MuggleMath (Li et al., 2024c)	ACL'24	English	37,365	P	E H	Discriminative	Open	S
OlympiadBench (He et al., 2024a) ★	ACL'24	English/Chinese	8,476	S	H C	Generative	Closed/Open/Math	S
MathBench (Liu et al., 2024b)	ACL Findings'24	English/Chinese	3,709	P S	E M H U	Generative	Closed/Open/Math	S
GeoEval (Zhang et al., 2024d) ★	ACL Findings'24	English	5,050	P G	E M H	Discriminative	Closed/Open/Math	S
QRData (Liu et al., 2024d)	ACL Findings'24	English	411	S	U	Discriminative	Closed/Open/Math	S
EIC-Math (Li et al., 2024e)	ACL Findings'24	English	1,800	P	E M H	Discriminative	Closed/Open	S
Srivastava et al. (2024)	ACL Findings'24	English	-	P	H	Discriminative	Closed/Open	D O
CHAMP (Mao et al., 2024)	ACL Findings'24	English	270	S	H	Generative	Closed/Open	S
IMO-AG-30 (Trinh et al., 2024)	Nature'24	English	30	S	C	Discriminative	Closed	P
PutnamBench (Tsoukalas et al.)	NeurIPS'24	English	1,697	S	C	Generative	Closed	S P
MATH-Vision (Wang et al., 2024a) ★	NeurIPS'24	English	3,040	S	E M H U	Discriminative	Closed/Open	S
CARP (Zhang et al., 2024a)	NeurIPS'24	Chinese	4,886	S	C	Discriminative	Closed	S
SMART-840 (Cherian et al., 2024) ★	NeurIPS'24	English	840	S	E M H C	Discriminative	Closed/Open	S
OpenMathInstruct-1 (Toshniwal et al., 2024b)	NeurIPS'24	English	1,800,000	P	E M H C	Generative	Closed/Open/Math	S
DidolKar et al. (2024)	NeurIPS'24	English	8,600	P	E	Discriminative	Closed	S O
Putnam-AXIOM (Gulati et al., 2024)	NeurIPS Workshop'24	English	236	S	C	Discriminative	Closed/Open/Math	S
Scibench (Wang et al., 2023b) ★	ICML'24	English	869	S	U	Discriminative	Closed/Open	S
GeomVerse (Kazemi et al., 2023) ★	ICML Workshop'24	English	1,000	G	U	Discriminative	Closed	S
MathVista (Lu et al., 2023) ★	ICLR'24	English	6,141	S P	E M H U	Discriminative	Closed/Open	S
MMMU _{math} (Yue et al., 2024a) ★	CVPR'24	English	540	S	U	Discriminative	Closed/Open	S
MathVerse (Zhang et al., 2024f) ★	ECCV'24	English	2,612	S P	H	Generative	Closed/Open	S
Mathador-LM (Kurtic et al., 2024)	EMNLP'24	English	-	G	E	Both	Closed/Open	S D
MM-MATH (Sun et al., 2024a) ★	EMNLP Findings'24	English	5,929	S	M H	Discriminative	Closed/Open	S D
Scieval (Sun et al., 2024b)	AAAI'24	English	15,901	S P	H	Both	Closed/Open	S
ArqMATH (Satpute et al., 2024)	SIGIR'24	English	450	P	U	Generative	Closed/Open/Math	S
IsoBench (Fu et al., 2024a) ★	COLM'24	English	1,887	S	E M H U	Discriminative	Closed/Open	S
MMMU-Pro _{math} (Yue et al., 2024b) ★	arXiv'24	English	60	S	U	Discriminative	Closed/Open	S
MathOdyssey (Fang et al., 2024)	arXiv'24	English	387	S	H U C	Both	Closed/Open/Math	S
MathScape (Zhou et al., 2024b) ★	arXiv'24	Chinese	1,325	S	E M H C	Generative	Closed/Open	S
U-Math (Chernyshev et al., 2024) ★	arXiv'24	English	1,100	S	U	Discriminative	Closed/Open/Math	S D
MathHay (Wang et al., 2024b)	arXiv'24	English	673	S P	H	Both	Closed/Open	S
ErrorRador (Yan et al., 2024a) ★	arXiv'24	English	2,500	S	E M H	Discriminative	Closed/Open	D
FaultyMath (Rahman et al., 2024) ★	arXiv'24	English	363	G	E M H	Discriminative	Closed/Open/Math	D
MathChat (Liang et al., 2024c)	arXiv'24	English	1,319	P	E	Both	Closed/Open/Math	S D O
E-GSM (Xu et al., 2024e)	arXiv'24	Chinese	4,500	P	E	Both	Closed/Open/Math	S O
Tangram (Tang et al., 2024a) ★	arXiv'24	English	4,320	S	E M H C	Discriminative	Closed/Open	S
CMM-Math (Liu et al., 2024c) ★	arXiv'24	Chinese	28,069	S	E M H	Both	Closed/Open/Math	S
CMMaTH (Li et al., 2024i) ★	arXiv'24	English/Chinese	23,856	S	E M H	Both	Closed/Open/Math	S
EAGLE (Li et al., 2024h) ★	arXiv'24	English	170,000	P	E M H	Discriminative	Closed/Open/Math	S
VisAidMath (Ma et al., 2024) ★	arXiv'24	English	1,200	S	M H C	Discriminative	Closed/Open	S
AutoGeo (Huang et al., 2024d) ★	arXiv'24	English	100,000	S	E M H U	Both	Closed/Open	S
NTKEval (Guo et al., 2024a)	arXiv'24	English	1,860	P G	H	Discriminative	Open	S
Mamo (Huang et al., 2024b)	arXiv'24	English	1,209	S G	U	Generative	Closed/Open/Math	S
RoMath (Cosma et al., 2024)	arXiv'24	Romanian	70,000	S	M H C	Discriminative	Closed/Open/Math	S
MaTT (Davoodi et al., 2024)	arXiv'24	English	1,958	S	U	Discriminative	Closed/Open	S
Li et al. (2024a)	arXiv'24	English	15,000	P	E M H	Generative	Closed/Open/Math	S
PolyMATH (Gupta et al., 2024) ★	arXiv'24	English	5,000	S	M H U	Discriminative	Closed/Open	S
SuperCLUE-Math6 (Xu et al., 2024b)	arXiv'24	English/Chinese	2,144	S	E	Generative	Closed/Open	S
TheoremQA (Chen et al., 2023)	EMNLP'23	English	800	S	U	Discriminative	Closed/Open	S
LILA (Mishra et al., 2022)	EMNLP'22	English	133,815	P	H	Discriminative	Closed	S
GeoQA (Chen et al., 2021) ★	ACL'21	Chinese	4,998	S	M	Discriminative	Open	S
MATH (Hendrycks et al., 2021)	NeurIPS'21	English	12,500	S	C	Discriminative	Closed	S

Table 3: **Overview of LLM-based benchmarks for mathematical reasoning.** ★ refers to those designed to evaluate the multimodal mathematical setting. Different colors indicate different types for the following columns: **Source:** S = Self-Sourced, P = Collected from Public Dataset, G = Generated by LLM **Level:** E = Elementary, M = Middle School, H = High School, U = University, C = Competition, H = Hybrid **Task:** S = Problem-Solving, D = Error Detection, P = Proving, O = Others

the incorporation of three-dimensional spatial reasoning and the need to analyze geometric properties of solids rather than flat, planar relationships. These tasks challenge students to bridge visual understanding with formulas involving multiple dimensions.

C.3 Multimodal Diagram Setting

In the Multimodal Diagram Setting, the problems revolve around interpreting visual data presented in the form of tables, charts, or diagrams. These tasks require students to extract numerical or cate-

gorical information and perform basic operations, such as addition, comparison, or selection. The visual components often include neatly organized tables, bar charts, or pie graphs, where the information is clearly labeled for accessibility. Unlike geometry-based problems, which require spatial reasoning, diagram settings focus on numerical literacy and the ability to synthesize information from structured visual data. This type highlights the integration of simple arithmetic and the comprehension of organized visual representations.

Math (LLMs)	Organization	Release Date	Publication	Language	Parameter Size	Evaluation Benchmarks	Open Source
GPT-4 (Pola and Sutskever, 2021)	OpenAI	Sep 2020	-	English	160M/400M/700M	-	✓
Hypertree Proof Search (Lample et al., 2022)	Meta	Nov 2022	NeurIPS '22	English	-	miniF2F/Metamath	✓
Minerva (Lewkowycz et al., 2022)	Google	Jun 2022	NeurIPS '22	English	8B/62B/540B	MATH/MMLU-STEM/GSM8k	✓
JiuZhang 1.0 (Zhao et al., 2022)	RUC & iFLYTEK	Jun 2022	KDD '22	English	1.45M	-	✓
GABR-Math-Abel (Chen et al., 2023)	Shanghai Jiaotong University	2023	-	English	7B/13B/70B	GSM8k/MATH/MMLU/SVAMP/SCQ5K-English/MathQA	✓
JiuZhang 2.0 (Zhao et al., 2023)	RUC & iFLYTEK	Jun 2023	KDD ADS '23	English	-	JCAI/IBAG (MathBERT/DART/JiuZhang)	✓
KwaiYiMath (Fu et al., 2023)	Kaishou	Jan 2023	-	English/Chinese	13B	GSM8k/MATH/KMath	✓
MathCoder (Wang et al., 2023a)	CUHK	Jan 2023	ICLR '24	English	7B/13B	GSM8k/MATH	✓
Llemma (Azerbayev et al., 2023)	Princeton University & Eleuther AI	Jan 2023	-	English	7B/34B	MATH/GSM8k/MMLU-STEM/AT/OCW/Course	✓
Skywork-13B-Math (Zeng et al., 2024) ★	SkyworkAI	Jan 2023	-	English	7B/13B	GSM8k/MATH/MATH	✓
MathGPT (TALEducation, 2023) ★	TAL Education Group	Aug 2023	-	English/Chinese	130B	CEval-Math/AGIEval-Math/APESK/CMMLU-Math/GAOKAO-Math/Math401	✓
WizardMath (Luo et al., 2023)	Microsoft	Aug 2023	ICLR '25	English	7B/70B	GSM8k/MATH	✓
MAMmo/THI (Yue et al., 2023)	UWaterloo	Sep 2023	ICLR '24	English	7B/13B/70B	GSM/MATH/MMLU-STEM/AQua/NumGLUE	✓
MathGLM (Yang et al., 2023b)	Tsinghua & Zhipu AI	Sep 2023	-	English	10M/100M/500M/2B/Arith. & 235M/6B/10B (MWP)	BIG-bench/ Ape210K	✓
MetaMath (Yu et al., 2023)	Cambridge & Huawei	Sep 2023	-	English	7B/13B/70B	GSM8k/MATH	✓
DeepSeekMath (Shao et al., 2024)	DeepSeek AI	Jan 2024	-	English	7B	GSM8k/MATH/OCW/SAT/MMLU-STEM/CMATH/Gaokao-Math/Cloze/Gaokao-MathQA	✓
InternLM2.5-StepProver (Wu et al., 2024c)	Shanghai AI Lab	Jan 2024	-	English/Chinese	7B	miniF2F/Lean-WorkBook-Plus/ProofNet/Putnam	✓
ChatGLM-Math (Xu et al., 2024f)	Zhipu AI	Apr 2024	-	English/Chinese	32B	MathUserEval/Ape210K/CMATH/GSM8k/MATH/Hungarian	✓
Rho-Math (Lin et al., 2024)	Microsoft	Apr 2024	-	English	1B/7B	GSM8k/MATH/MMLU-STEM/SAT/SVAMP/ASDiv/MAWPS/TAB/MQA	✓
DeepSeekProver-V1 (Xin et al., 2024b)	DeepSeek AI	May 2024	-	English	7B	miniF2F/FIMO	✓
InternLM2-Math (Wu et al., 2024c)	Shanghai AI Lab	May 2024	-	English/Chinese	1.8B/7B/20B/8x22B	MiniF2F-test/MATH/MATH-Python/GSM8k/MathBench-A/Hangary/	✓
JiuZhang 3.0 (Zhao et al., 2024a)	RUC & iFLYTEK	May 2024	NeurIPS '24	English	7B/8B	GSM8k/MATH/G-Head/SVAMP/MAWPS/ASDiv/TAB/MWP	✓
MAMmo/TH2 (Yue et al., 2024c)	UWaterloo	May 2024	-	English	7B/8B	TheoremQA/MATH/GSM8k/GPQA/MMLU-STEM/BHH	✓
Math-LLaVA (Shi et al., 2024)	NUS	Jun 2024	EMNLP Finding '24	English	13B	MMU/MATH-V/MathVista	✓
Mathstral (MistralAI, 2024)	Mistral AI	Jul 2024	-	English	7B	MATH/GSM8k/GRE/Math/AMC2023/AIME2024/MathOdyssey	✓
DeepSeek-Prover-V1.5 (Xin et al., 2024a)	DeepSeek AI	Aug 2024	-	English	7B	miniF2F-test/ProofNet	✓
Qwen2-Math (Qwen, 2024)	Alibaba	Aug 2024	-	English/Chinese	1.5B/7B/72B	GSM8k/MATH/MMLU-STEM/CMATH/GaokaoMath Cloze/Gaokao Math QA	✓
Qwen2-Math-Instruct (Qwen, 2024)	Alibaba	Aug 2024	-	English/Chinese	1.5B/7B/72B	GSM8k/MATH/Mimosa Math/GaoKao2023 En/Olympiad Bench/College Math/MMLU STEM/Gaokao/CMATH/CNMiddle School 24/AIME24/AMC23	✓
MathGLM-Vision (Yang et al., 2024b) ★	Tsinghua & Zhipu AI	Sep 2024	-	English	9B/19B/32B	MathVista/MathVista(GPS)/MathVerse/Math-Vision/MMU/MathVL	✓
Math-LLM (Liu et al., 2024c) ★	East China Normal University	Sep 2024	-	Chinese	8.26B/7B/72B	CMM-Math/MathVista/Math-V	✓
Qwen2.5-Math (Yang et al., 2024a)	Alibaba	Sep 2024	-	English/Chinese	1.5B/7B/72B	GSM8k/MATH/MMLU-STEM/CMATH/Gaokao Math	✓
Xwin-1M (Ni et al., 2024)	Microsoft	May 2024	-	English	7B/13B/70B	GSM8k/MATH	✓
MathCoder2 (Liu et al., 2024c)	CUHK	Nov 2024	ICLR '25	English	7B	GSM8k/MATH/SAT-Math/OCW/MMLU-Math	✓
math-specialized Gemini 1.5 Pro ★	Google	Not launched yet	-	English	-	MATH/AIME2024/MathOdyssey/HiddenMath/IMO Bench	-
k0-math (MoonshotAI, 2024)	Moonshot AI	Nov 2024	-	English/Chinese	-	KAOYAN/MATH/AIME/OMNI-MATH/GAOKAO/ZHONGKAO	-
Duolingo Math (Duolingo, 2024)	Duolingo	2024	-	English	-	-	-
Khanmigo (KhanAcademy, 2024)	Khan Academy	2024	-	English	-	-	-
Squirrel LAM (SquirrelAILearning, 2024) ★	Squirrel AI Learning	2024	-	Chinese	-	-	-

Table 4: **Overview of math-specific LLMs (sort by release date).** ★ refers to those designed to support the multimodal mathematical setting.

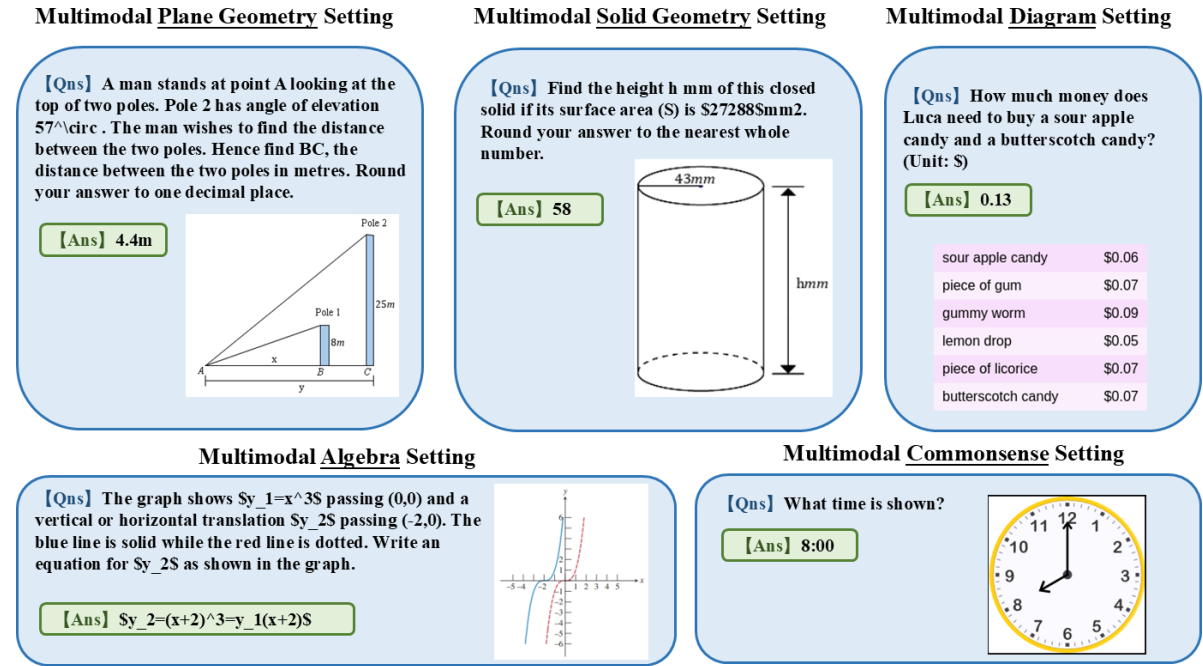


Figure 5: The illustration of diverse multimodal mathematical settings.

C.4 Multimodal Algebra Setting

The Multimodal Algebra Setting introduces problems that combine graphical representations and algebraic reasoning. These tasks often involve interpreting visual graphs, identifying equations, or understanding transformations such as translations or reflections. The visuals typically feature coordinate graphs with curves or lines, where solid and dotted lines may represent different functions or changes. Students are required to connect the

visual graph to algebraic expressions, such as equations or transformations of functions. This type of question emphasizes the interplay between visual understanding (graph) and symbolic representation (algebra), making it distinct from purely numerical or geometric settings.

C.5 Multimodal Commonsense Setting

The Multimodal Commonsense Setting is characterized by problems that involve interpreting everyday visuals and applying logical reasoning. These ques-

tions present familiar objects, such as clocks, calendars, or real-world scenarios, where students must analyze the visual information to derive straightforward answers. Visually, these tasks feature clear and relatable imagery, like an analog clock with its hands pointing to a specific time. Unlike other types, commonsense settings rely less on abstract mathematical reasoning and more on practical interpretation of everyday visual cues. This setting highlights how mathematical understanding can intersect with routine, real-world observations.

C.6 Summary

In summary, the key differences among these types stem from their visual focus and cognitive demands. While plane and solid geometry emphasize spatial reasoning in 2D and 3D, respectively, diagram settings target numerical literacy through organized data. Algebra settings merge visual graphs with algebraic transformations, and commonsense settings leverage real-world visuals requiring practical logic. Each type uniquely integrates multimodal elements to challenge students across different mathematical skills.

D Details of Metrics

D.1 Discriminative Metrics

Discriminative tasks refer to evaluation processes where the outputs are typically binary, such as "Yes" or "No". These tasks often include multiple-choice questions, fill-in-the-blank problems, or judgment assessments. The evaluation metrics focus on LLM's accuracy in specific task types and its ability to control biases.

Accuracy (ACC): It measures the proportion of correctly predicted outcomes. The value should be as high as possible.

$$ACC = \frac{\sum_{1,m} x_i}{\sum_{1,n} y_j}$$

Where x_i represents the correct output for the i -th instance, y_j represents the j -th instance, m is the number of the correct instances and n is the number of the total instances.

Exact match: It evaluates the congruence between the answers generated by LLM and the correct ones. Specifically, in cases where the answer produced LLM coincides with the reference answer, a score of 1 point will be assigned. Conversely, if there is any discrepancy between them except for bias, a score of 0 point will be given.

F_1 score: It combines two crucial aspects, namely precision and recall, in order to comprehensively assess the accuracy of LLM. It is calculated as :

$$F_1 = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

The value of the F_1 score ranges from 0 to 1. A higher value of the F_1 score indicates better overall performance of LLM in terms of both precision and recall.

Macro- F_1 score: It calculates the F_1 score for each category separately and then takes the average of the F_1 scores of all categories, so as to obtain the overall performance of LLM on all categories.

Round-r accuracy: It is the proportion of correct answers given by a model on the question set Q_r in round r . It is calculated as follows:

$$ACC_r(M) = \frac{\sum_{q \in Q_r} I[M(q) = g_t(q)]}{|Q_r|}$$

Here, $ACC_r(M)$ represents the accuracy of LLM M on question set Q_r in round r . I is an indicator function. When the answer $M(q)$ given by M for question q is consistent with the true answer $g_t(q)$ of the question, the value of I is 1; otherwise, it is 0. The symbol $\sum_{q \in Q_r}$ means summing over all questions in question set Q_r . $|Q_r|$ indicates the number of questions in question set Q_r .

ACC_{step} : It is used to evaluate LLM's ability to identify the first step where an error occurs. The accuracy for identifying the first erroneous step is calculated as follows:

$$ACC_{step} = \frac{1}{N} \sum_{i=1}^N \mathbb{I}(S_{step,i} = G_{step,i})$$

Here, N is the total number of samples. For the i -th sample, $S_{step,i}$ is the predicted step where the error occurs, and $G_{step,i}$ is the ground truth label for the first erroneous step. The indicator function $\mathbb{I}(\cdot)$ returns 1 if the predicted step matches the ground truth and 0 otherwise.

ACC_{cate} : It is for assessing LLM's performance in categorizing the type of error. The accuracy for error categorization is defined by

$$ACC_{cate} = \frac{1}{N} \sum_{i=1}^N \mathbb{I}(C_{error,i} = G_{error,i})$$

Here, N is the total number of samples. For the i -th sample, $C_{error,i}$ is the predicted error category, and $G_{error,i}$ is the ground truth label for the error category. The indicator function $\mathbb{I}(\cdot)$ has the same

meaning as in the previous metric, returning 1 if the predicted error category matches the ground truth and 0 otherwise.

The skill success rate: It measures the proportion of a model correctly applying major skills in problem-solving. It’s calculated by analyzing test questions and determining correct use of major skills, then finding the ratio to total questions. For example, in triangle area calculation, checking use of the area formula. Similarly, **the secondary skill success rate** focuses on the proportion of correct application of secondary skills like understanding graphic properties and unit conversion, calculated by analyzing problem-solving and finding the ratio to total questions.

The False Positive Rate (FPR): It is the proportion of cases where the evaluation LLM misjudges an incorrect answer as a correct one. A low FPR indicates that LLM rarely misjudges incorrect student answers as correct.

The False Negative Rate (FNR): It is the proportion of cases where the evaluation LLM misjudges a correct answer as an incorrect one. A low FNR indicates that LLM is relatively accurate in correctly determining whether a student’s answer is correct.

Mean Squared Error (MSE): It is a metric that measures the average of the squares of the differences between the LLM’s predicted values and the actual true values. It is calculated as:

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

Here, n represents the number of samples. For the i -th sample, y_i is the true value and $\hat{y}_i$ is the predicted value by LLM. The summation symbol $\sum_{i=1}^n$ means summing up the squared differences for all n samples. Dividing by n gives the average squared difference, which is the MSE. MSE should be as low as possible.

Average-Case Accuracy (A_{avg}): This metric evaluates the average accuracy of LLM across all variants of a seed question. It is calculated as the proportion of correct answers across all variants and seed questions. The formula is:

$$A_{avg} = \frac{1}{N} \sum_{i=1}^N \frac{1}{M} \sum_{j=1}^M I[\text{Ans}(i, j) = \text{GT}(i, j)]$$

where N is the total number of seed questions, M is the number of variants per seed question, and $I[\text{Ans}(i, j) = \text{GT}(i, j)]$ checks if the answer matches the ground truth.

Worst-Case Accuracy (A_{wst}): This evaluates the worst-case performance by considering the minimum accuracy across all variants of a seed question. It reflects the robustness of LLM against challenging variations. The formula is:

$$A_{wst} = \frac{1}{N} \sum_{i=1}^N \min_{j \in [1, M]} I[\text{Ans}(i, j) = \text{GT}(i, j)]$$

D.2 Generative Metrics

Generative tasks involve evaluating the content generated by LLM, typically encompassing free-form answers and responses to open-ended questions. These tasks focus primarily on assessing the extent of hallucinations in the generated content, especially when the content is not faithful to the given images. Evaluating generative tasks often requires more complex metrics, such as CHAIR and Faithscore, which measure hallucinations across different categories, including objects, attributes, and relationships within the generated content. These metrics provide a nuanced understanding of the fidelity and reliability of MLLMs in producing content aligned with the visual and textual inputs.

Reasoning Robustness (RR): This metric measures the relative robustness of LLM by comparing the worst-case performance to the average-case performance. The formula is:

$$RR = \frac{A_{wst}}{A_{avg}}$$

Repetition Consistency (RC): This evaluates the consistency of LLM’s responses across repeated queries for the same question variant. It helps distinguish between variability due to randomness and systematic errors. The formula is:

$$RC(i, j) = \frac{1}{K} \sum_{k=1}^K I[\text{Ans}_k(i, j) = \text{Ans}(i, j)]$$

where K is the number of repetitions.

OpenCompass Scoring: It is a comprehensive evaluation framework that leverages the OpenCompass platform to assess the generative capabilities of LLM across multiple dimensions. Perplexity (PPL) evaluates the naturalness and fluency of generated text, with lower scores indicating greater model confidence and the ability to produce contextually coherent sequences. Simultaneously, CircularEval assesses the robustness and consistency of LLM in multiple-choice scenarios by evaluating its performance across N random permutations of

the options in an N -option question. A question is deemed correctly answered only if LLM provides the correct response for all permutations, highlighting its ability to handle randomized inputs reliably.

Bilingual Evaluation Understudy (BLEU): It evaluates the quality of text generation by measuring n-gram overlap between generated and reference texts, focusing on precision and brevity. Its formula is:

$$\text{BLEU} = \text{BP} \cdot \exp \left(\sum_{n=1}^N w_n \log p_n \right)$$

where BP is the brevity penalty, calculated as 1 if $c > r$, or $\exp(1 - r/c)$ if $c \leq r$, with c and r representing the lengths of the generated and reference texts, respectively. w_n denotes n-gram weights (typically uniform), and p_n is the precision of n-grams of size n . BLEU scores range from 0 to 1 (often expressed as percentages, 0-100%), with higher scores indicating greater similarity between the generated and reference texts.

Recall-Oriented Understudy for Gisting Evaluation-L (ROUGE-L): It evaluates the quality of generated text by measuring its similarity to reference text, focusing on sequence alignment and structural consistency through the Longest Common Subsequence (LCS). It calculates recall as the proportion of the LCS length relative to the reference text length. The formula of recall is:

$$R = \frac{\text{LCS}(\text{Generated}, \text{Reference})}{\text{Length}(\text{Reference})}$$

It also calculates precision as the proportion of the LCS length relative to the generated text length. The formula is:

$$P = \frac{\text{LCS}(\text{Generated}, \text{Reference})}{\text{Length}(\text{Generated})}$$

The F_1 score is a harmonic mean of precision and recall, expressed as:

$$F_1 = \frac{(1 + \beta^2) \cdot P \cdot R}{\beta^2 \cdot P + R}$$

where β (commonly set to 1) controls the weighting of recall and precision. ROUGE-L scores range from 0 to 1, with higher scores indicating greater similarity between the generated and reference texts.

Consensus-based Image Description Evaluation (CIDEr): It is designed for image description tasks, measuring the semantic relevance of

generated descriptions by calculating the TF-IDF weighted n-gram similarity with reference descriptions. The formula is:

$$\text{CIDEr}_n(c_i, S_i) = \frac{1}{m} \sum_{j=1}^m \frac{g^n(c_i) \cdot g^n(s_{ij})}{\|g^n(c_i)\| \cdot \|g^n(s_{ij})\|}$$

$$\text{CIDEr}(c_i, S_i) = \sum_{n=1}^N w_n \text{CIDEr}_n(c_i, S_i)$$

Here, c_i is the candidate description, $S_i = \{s_{i1}, s_{i2}, \dots, s_{im}\}$ is the set of reference descriptions, and m is the number of references. $g^n(c_i)$ and $g^n(s_{ij})$ are the TF-IDF weighted n-gram vectors for the candidate and reference descriptions, with $\|g^n(c_i)\|$ and $\|g^n(s_{ij})\|$ being their magnitudes. w_n is the weight for n-grams of different lengths, usually $w_n = 1/N$, where N is the maximum n-gram length. Scores range from 0 to 10, with higher scores indicating stronger alignment between candidate and reference descriptions.

Mathematical Symbol Similarity: This metric measures the similarity between the correct steps in a reasoning process and the steps generated by LLM, using symbolic computation software to perform the evaluation.

GPT Scoring: This metric evaluates the generated content based on scores assigned by GPT or other language models, focusing on the linguistic coherence and logical consistency of the text.

Context Length Generalization Efficacy (CoLeG-E): It is a metric used to measure LLM's consistency in answering variations of the same question across different context lengths. It is defined as:

$$\text{CoLeG-E}(M) = \frac{\sum_{q \in Q_R} [\bigwedge_{r=1}^R I[M(q^r) = gt(q^r)]]}{|Q_R|}$$

where Q_R represents the set of all questions under evaluation, and q^r refers to the r -th variation of a question q , corresponding to a specific context length. $M(q^r)$ is LLM's predicted answer for the r -th variation, while $gt(q^r)$ denotes the ground truth answer. The indicator function $I[\cdot]$ equals 1 if LLM's answer matches the ground truth, and 0 otherwise. The logical AND operator $\bigwedge_{r=1}^R$ ensures that the model must answer all variations of a question correctly for that question to be considered correctly answered.

Context Length Generalization Robustness (CoLeG-R): It measures LLM's robustness to context length expansion by quantifying the relative

drop in accuracy from initial to extended questions. It is defined as:

$$CoLeG-R(M) = 1 - \frac{ACC_0(M) - ACC_R(M)}{ACC_0(M)}$$

Here, $ACC_0(M)$ is the LLM’s accuracy on the initial set of shorter-context questions Q_0 , and $ACC_R(M)$ is its accuracy on the extended longer-context questions Q_R . Higher CoLeG-R values indicate better robustness, with less performance degradation across context lengths.

Performance Drop Rate (PDR): This metric measures the relative decline in model performance when transitioning from the original dataset to the perturbed dataset. It is defined as:

$$PDR = 1 - \frac{\sum_{(x,y) \in D_a} I[LLM(x), y] / |D_a|}{\sum_{(x,y) \in D} I[LLM(x), y] / |D|}$$

where D is the original dataset and D_a is the perturbed dataset. $I[LLM(x), y]$ is an indicator function that checks if the LLM’s output matches the ground truth y .

Accurately Solved Pairs (ASP): ASP measures the percentage of seed questions and their perturbed variations that are both correctly answered by LLM. It is defined as:

$$ASP = \frac{\sum_{x,y,x',y'} I[LLM(x), y] \cdot I[LLM(x'), y']}{N \cdot |D|}$$

where x and x' are a seed question and its variation, respectively. N is the number of perturbations per question. $|D|$ is the total number of seed questions.

Mean Average Precision (mAP): It is a metric that evaluates LLM’s ability to rank relevant answers higher in its output list for a given query. It is defined as:

$$mAP = \frac{1}{|Q|} \sum_{q \in Q} AP(q)$$

$$AP(q) = \frac{1}{m} \sum_{k=1}^m P(k)$$

$$P(k) = \frac{\# \text{ relevant ans retrieved up to position } k}{k}$$

Here, Q represents the set of all queries in the dataset. $AP(q)$ is the Average Precision for query q , calculated as the mean of the precision values $P(k)$ at ranks where relevant answers appear. $P(k)$ is the precision at rank k , representing the proportion of relevant answers retrieved up to position k .

m is the total number of relevant answers for query q .

Training Set Coverage (TSC): It measures how effectively LLM has learned to generate correct solutions for tasks similar to those in its training set. TSC is particularly useful in cross-domain or cross-modal tasks, where it assesses LLM’s ability to generalize learned patterns to problems aligned with its training data. Higher TSC scores indicate better learning and consistency, while lower scores suggest insufficient training or overfitting.

Pass@N: This metric measures the likelihood of LLM generating at least one correct solution within N attempts for a given problem. Formally:

$$Pass@N = \mathbb{E}_{\text{Problems}}[\min(c, 1)]$$

where c represents the number of correct answers out of N responses. A higher Pass@N indicates a greater chance of producing a correct answer in multiple attempts, reflecting LLM’s potential capability.

PassRatio@N: This metric calculates the proportion of correct answers among N generated responses for a given problem. It is defined as:

$$PassRatio@N = \mathbb{E}_{\text{Problems}} \left[\frac{c}{N} \right]$$

where c is the count of correct answers. This metric reflects LLM’s stability in consistently generating correct answers. It can be considered analogous to Pass@1 but offers reduced variance.

E Summary of Methods

Table 5 summarizes the LLM-based methods for mathematical reasoning.

F More Details of Challenges

F.1 Discussion of Data Bottlenecks

We dive into the three bottlenecks of multimodal mathematical datasets as follows.

❶ Bottleneck in Data Quality:

1. **Labeling Noise and Modality Alignment:** Multimodal math problems often involve complex associations between text, formulas, and charts. Mismatches between text descriptions and images/formulas (e.g., incorrect axis labels, contradictions between geometry figures and problem statements) can severely impair the model’s ability to understand cross-modal relationships.

Methods	Venue	Evaluated Math Dataset(s)	Task(s)	Scope(s)	LLM as Enhancer	LLM as Reasoner	LLM as Planner
MAVIS (Zhang et al., 2024g) ★	ICLR'25	MathVerse/GeoQA/MathVista/MMMU/MathVision	S	M	✓	✓	
TVM (Lee et al., 2024)	ICLR'25	GSM8K/MATH	S	A		✓	
MathCoder2 (Lu et al., 2024c)	ICLR'25	GSM8K/MATH/SAT-Math/OCW/MMLU-Math	S P	M	✓	✓	
Xiong et al. (2024)	ICLR'25	GSM8K/MATH	S	A		✓	✓
TSMC (Feng et al., 2024)	ICLR'25	GSM8K/MATH500	S P	A		✓	
AlphaGeometry (Trinh et al., 2024) ★	Nature'24	IMO-AG-30	S P	G	✓	✓	
Masked Thought (Chen et al., 2024a)	ACL'24	GSM8K/MATH/GSM8K-RFT/MetaMathQA/MathInstruct	S	A	✓	✓	
MathGenie (Lu et al., 2024b)	ACL'24	GSM8K/MATH/SVAMP/Simuleq/Mathematics	S	A	✓	✓	
MATH-SHEPHERD (Wang et al., 2024c)	ACL'24	GSM8K/MATH	S P	A	✓	✓	
SEGO (Zhao et al., 2024)	ACL'24	GSM8K/MATH	S P	A	✓	✓	
Deng et al. (2023)	ACL Workshop'24	GSM8K/SVAMP/MultiArith/MathQA/CSQA	S P	A		✓	
MathCoder (Wang et al., 2023a)	ICLR'24	GSM8K/MATH	S P	A		✓	
ToRA (Gou et al., 2023)	ICLR'24	GSM8K/MATH	S P	A		✓	✓
Visual Sketchpad (Hu et al., 2024) ★	NeurIPS'24	Geometry3K/ IsoBench	S	G		✓	✓
JiuZhang 3.0 (Zhou et al., 2024a)	NeurIPS'24	GSM8K/MATH/SVAMP/ASDiv/MAWPS/CARP	S P	A	✓	✓	
Minimo (Poesia et al., 2024)	NeurIPS'24	MATH/GSM8K/College/DM/Olympiad/Theorem	S P	A	✓	✓	
DART-Math (Tong et al., 2024)	NeurIPS'24	GSM8K/MATH	S P	A	✓	✓	
Li et al. (2024f)	NeurIPS'24	MATH	S	A		✓	
MACM (Lei et al., 2024)	NeurIPS'24	MATH	S	A		✓	
Sinha et al. (2024) ★	NeurIPS Workshop'24	IMO-AG-30	S P	G		✓	
SBIRAG (Dixit and Oates, 2024)	NeurIPS Workshop'24	GSM8K	S	A		✓	
MathScale (Tang et al., 2024b)	ICML'24	-	S	A		✓	
VerityMath (Han et al., 2023)	ICML Workshop'24	GSM8K	S	A	✓	✓	
RefAug (Zhang et al., 2024j)	EMNLP'24	GSM8K/MATH/Mathematics/MAWPS/SVAMP/MMLU-Math/SAT-Math/MathChat-FQA/MathChat-EC/MINI-Math	S P	A	✓	✓	
Math-LLaVA (Shi et al., 2024) ★	EMNLP Findings'24	MathVista/Math-V	S P	M	✓	✓	
COPRA (Thakur et al., 2024)	COLM'24	miniF2F-test	S	A		✓	✓
PRP (Wu et al., 2024b)	AAAI'24	MAWPS/ASDivA/Math23k/SVAMP/Un-biasedMWP	S	A		✓	
PERC (Jin et al., 2024)	L@S'24	PERC	S	A		✓	
Math-PUMA (Zhuang et al., 2024) ★	arXiv'24	MathVerse/MathVista/WE-MATH	S P	M		✓	
MultiMath (Peng et al., 2024) ★	arXiv'24	MathVista/MathVerse/MultiMath-300K	S P	M		✓	
MathAttack (Zhou et al., 2024e)	arXiv'24	GSM8K/MultiArith	S	A		✓	
MinT (Liang et al., 2023b)	arXiv'24	GSM8K/MathQA/CM17k/Ape210k	S	A		✓	
DotaMath (Li et al., 2024b)	arXiv'24	GSM8K/MATH/Mathematics/SVAMP/TabMWP/ASDiv	S	A	✓	✓	
DFE-GPS (Zhang et al., 2024i)	arXiv'24	FORMALGEO7k	S	A		✓	
PGPSNet-v2 (Zhang et al., 2024e) ★	arXiv'24	Geometry3K/PGPS9K	S	G B	✓	✓	
LLaMA-Berry (Zhang et al., 2024b)	arXiv'24	GSM8K/MATH/GaoKao2023En/OlympiadBench/College Math/MMLU STEM	S P	A		✓	
Skywork-Math (Zeng et al., 2024) ★	arXiv'24	GSM8K/MATH	S P	A	✓	✓	
SLaM (Yu et al., 2024a)	arXiv'24	GSM8K/CMATH	S P	A		✓	
InternLM-Math (Ying et al., 2024)	arXiv'24	GSM8K/MATH	S P	A		✓	
MathGLM-Vision (Yang et al., 2024b) ★	arXiv'24	MathVista/MathVerse/MathVision	S P	M	✓	✓	
Qwen2.5-Math (Yang et al., 2024a) ★	arXiv'24	GSM8K/MATH/MMLU-STEM/CMATH/GaoKao-Math-Cloze/GaoKao-Math-QA	S P	A	✓	✓	
S3c-Math (Yan et al., 2024c)	arXiv'24	GSM8K/MATH/SVAMP/Mathematics	S P	A	✓	✓	
SIRP (Wu et al., 2024a)	arXiv'24	CSQA/GSM8K/MATH/MBPP	S P	A		✓	
AIPS (Wei et al., 2024)	arXiv'24	MO-INT-20	S	G		✓	
DeepSeekMath (Shao et al., 2024)	arXiv'24	GSM8K/MATH/OCW/SAT/MMLU STEM/CMATH-Gaokao MathCloze/Gaokao MathQA	S	A		✓	
MMIQc (Liu et al., 2024a)	arXiv'24	MATH/MMIQc	S	A	✓	✓	
LANS (Li et al., 2023c) ★	arXiv'24	Geometry3K/PGPS9K	S	G B		✓	
VCAR (Jia et al., 2024) ★	arXiv'24	MathVista/MathVerse	S	M		✓	
KPDDS (Huang et al., 2024c)	arXiv'24	GSM8K/MATH/SVAMP/TabMWP/ASDiv/MAWPS	S	A	✓	✓	
HGR (Huang et al., 2024a) ★	arXiv'24	Geometry3K	S	G	✓	✓	
InFiMM-Math (Han et al., 2024) ★	arXiv'24	GSM8K/MMLU/MathVerse/We-Math	S	A	✓	✓	
CoSC (Han et al., 2024)	arXiv'24	GSM8K/MATH	S P	A		✓	
SICCV (Liang et al., 2024b)	arXiv'24	GSM8K/MATH500	S P	A		✓	
BEATS (Sun et al., 2024c)	arXiv'24	GSM8K/MATH/SVAMP/SimulEq/NumGLUE	S P	A		✓	
MindStar (Kang et al., 2024)	arXiv'24	GSM8K/MATH	S P	A		✓	
UMM (Zhang et al., 2024h)	arXiv'24	MMLU/GSM8K-COT/GSM8K-Coding/MATH-COT/MATH-Coding/HumanEval/InfiBench	S P	A		✓	
STIC (Deng et al., 2024) ★	arXiv'24	ScienceQA/TextVQA/ChartQA/LLaVA-Bench/MMBench/MM-Vet/MathVista	S	M		✓	
SPMWP (Zhang et al., 2023)	ACL'23	GSM8K	S	A		✓	
CoRe (Zhu et al., 2022)	ACL'23	GSM8K/ASDiv-A/SingleOp/Sinleq/MultiArith	S	A		✓	
TabMWP (Lu et al., 2022a)	ICLR'23	TabMWP	S	A D		✓	
Chameleon (Lu et al., 2024a) ★	NeurIPS'23	ScienceQA/TabMWP	S	A D		✓	
ATHENA (Kim et al., 2023)	EMNLP'23	MAWPS/ASDivA/Math23k/SVAMP/Un-biasedMWP	S P	A		✓	✓
UniMath (Liang et al., 2023a) ★	EMNLP'23	SVAMP/GeoQA/TabMWP/MathQA/UniGeo-Proving	S P	A D		✓	
Jiuzhang 2.0 (Zhao et al., 2023)	KDD'23	MCQ/BFQ/CAG/BAG/KPC/QRC/JCAG/JBAG	S	A		✓	
TCDP (Qin et al., 2023)	TNNLS'23	Math23k/CM17K	S	A		✓	
UniGeo (Chen et al., 2022) ★	EMNLP'22	GeoQA/UniGeo	S P	G	✓	✓	
LogicSolver (Yang et al., 2022)	EMNLP Findings'22	InterMWP/Math23K	S	A		✓	
Jiuzhang (Zhao et al., 2022)	KDD'22	KPC/QRC/QAM/SQR/QAR/MCQ/BFQ/CAG/BAG	S	A		✓	
MWP-BERT (Liang et al., 2021)	NAACL'22	Math23k/MathQA/Ape-210k	S	A		✓	
Inter-GPS (Lu et al., 2021) ★	ACL'21	Geometry3K/GEOS	S	G		✓	

Table 5: **Overview of LLM-based methods for mathematical reasoning.** ★ refers to those specifically designed to tackle the multimodal mathematical setting. Different colors indicate different types for the following columns: **Task:** **S**= Problem-Solving, **D**= Error Detection, **P**= Proving, **O**= Others **Scope:** **G**= Geometry, **A**= Algebra, **D**= Diagram, **M**= General Math

2. **Lack of Deep Annotation for Problem Solving Process:** Most datasets only provide final answers, lacking intermediate steps such as algebraic transformations or construction of geometric auxiliary lines, making it difficult for models to learn the mathematical thinking chain (Chain-of-Thought).

2 Bottleneck in Data Diversity:

1. **Limited Coverage of Problem Types and Scenarios:** Existing datasets are mostly focused on basic math areas (*e.g.*, algebraic equations, simple geometry) and insufficiently cover higher-level math (*e.g.*, topology, discrete mathematics) or real-world scenarios (*e.g.*, physics modeling, financial calculations).

2. **Monotony in Multimodal Combination Patterns:** Modal interactions are often simple con-

2266	catenations (e.g., text + static charts) without	F.2 Limited Domain Generalization in	2315
2267	dynamic interactions (e.g., scalable geometric	Multimodal Contexts	2316
2268	figures), or multi-step cross-modal reasoning	We further discuss the challenge of <i>limited domain</i>	2317
2269	(e.g., generating charts from text descriptions	<i>generalization</i> in multimodal contexts through the	2318
2270	and then solving problems).	perspective of the methodology paradigm.	2319
2271	③ Bottleneck in Data Scale:		
2272	1. High Cost of High-Quality Data Acquisition: Mathematical problems need to be de-	1. LLM as Enhancer: Generate mixed-domain	2320
2273	signed by experts and ensure multimodal con-	problems (e.g., combining algebraic equations	2321
2274	sistency, which leads to long production cy-	with geometric figures) to force the model	2322
2275	cles and high costs. Additionally, there is data	to learn cross-domain associations. Explic-	2323
2276	scarcity for long-tail problems (e.g., niche	itly add domain labels (e.g., "spatial reason-	2324
2277	branches of mathematics), which cannot be	ing" label for geometry problems) to guide	2325
2278	supplemented by scraping existing resources	the model in distinguishing domain-specific	2326
2279	(e.g., textbooks, online question banks).	features. The limitation of this paradigm is	2327
2280		that enhanced data may lack the real-world	2328
2281	2. Imbalance in Modal Data Volumes: Text	complexity of domain intersections.	2329
2282	data volumes far exceed those of image/sym-		
2283	bol modalities, leading to models' insuffi-	2. LLM as Reasoner: Fine-tune the model sep-	2330
2284	cient feature extraction capability for non-text	arately for different mathematical domains	2331
2285	modalities.	(e.g., algebra, geometry) to learn domain-	2332
2286		specific visual patterns (e.g., encoding geomet-	2333
2287	④ Based on recent trends in the latest works, we	ric properties in figures). Use domain-specific	2334
2288	further propose the following actionable sugges-	few-shot examples (e.g., providing figure-text	2335
2289	tions to address these dataset bottlenecks:	associations in geometry) to guide the model	2336
2290	1. Innovation in Data Generation Techniques:	in switching reasoning modes. The limitation	2337
2291	Combine formal mathematical engines (e.g.,	is that the model's capacity may be limited,	2338
2292	Lean, Coq) to generate verifiable reasoning	making it difficult to master multiple signifi-	2339
2293	steps, use programmatic rendering tools (e.g.,	cantly different domains simultaneously (e.g.,	2340
2294	TikZ, GeoGebra) to automatically generate	switching from algebraic symbol manipula-	2341
2295	precise charts, and design semi-automated an-	tion to geometric spatial reasoning).	2342
2296	notation pipelines that reduce manual labor		
2297	through large models generating drafts and	3. LLM as Planner: Based on the problem	2343
2298	experts refining them.	domain (e.g., detecting the "triangle" key-	2344
2299	2. Diversity Enhancement Strategies: Con-	word), call specialized tools (e.g., geomet-	2345
2300	struct interdisciplinary, cross-cultural bench-	ric theorem prover). For composite problems	2346
2301	mark datasets (e.g., math-physics cross-	(e.g., algebraic-geometry equations), coordi-	2347
2302	domain problems), utilize crowdsourcing plat-	nate symbolic computation tools (e.g., Mathe-	2348
2303	forms to collect real-world scenario problems,	matica) and graphical reasoning tools (e.g.,	2349
2304	and explore controllable data augmentation	GeoGebra). The limitation is that domain	2350
2305	techniques, such as rule-based problem defor-	boundary issues (e.g., math word problems	2351
2306	mation (e.g., modifying parameters or replac-	requiring commonsense reasoning) may fail	2352
2307	ing chart elements).	to route to the appropriate tools.	2353
2308	3. Scaling and Resource Integration: Encour-	F.3 Error Feedback Limitations in	2354
2309	age collaborative dataset creation within the	Multimodal Contexts	2355
2310	academic community (similar to ProofWiki),	We further discuss the challenge of <i>error feedback</i>	2356
2311	integrate existing educational resources (e.g.,	<i>limitations</i> in multimodal contexts through the per-	2357
2312	Khan Academy video-text analysis), and use	spective of the methodology paradigm.	2358
2313	synthetic data to fill long-tail gaps while im-		
2314	proving model robustness to synthetic noise	1. LLM as Enhancer: Inject cross-modal errors	2359
	through adversarial training.	(e.g., plot errors in function curves while the	2360
		text description is correct) to train the model	2361
		to detect contradictions. The limitation is that	2362

labeling error types is costly and it’s difficult to cover all long-tail errors.

2. **LLM as Reasoner:** Decompose reasoning into "computation-logic-conclusion" stages and cross-check text derivations with graphical information (*e.g.*, verify function extrema calculations using coordinates in the image). The limitation is that self-doubt relies on the model’s prior knowledge of error types, potentially missing rare error patterns in the training data.
3. **LLM as Planner:** Use OCR tools to extract symbols from figures and compare them with the text description to detect misunderstandings. The limitation is that tool invocation delays affect real-time performance, and some errors require manually defined detection rules.

F.4 How Test-Time Scaling Techniques Handle Other Challenges

We believe that test-time scaling techniques (Xu et al., 2025a; Li et al., 2025; Besta et al., 2025; Muennighoff et al., 2025; Chen et al., 2024c, 2025) can also help handle other challenges discussed in Section 4, especially the following three challenges.

❶ Insufficient Visual Reasoning:

1. **Enhancement of Multimodal Reasoning Chains:** During reasoning, generate multi-step visual-symbol joint inference paths. For example, using CoT prompts to guide the model in decomposing geometric shapes into angle, side length, and other symbolic constraints, and then calling a geometry solver to validate spatial relationships.
2. **Visual-Symbol Alignment Verification:** Use Best-of-N sampling to generate multiple candidate diagram parsing results and call external OCR tools or geometry validators (*e.g.*, GeoGebra) to detect visual-text consistency and filter out erroneous explanations.
3. **Limitations:** Parsing complex visual details (*e.g.*, topological structures) depends on the pretrained visual encoder’s capabilities. If the training data coverage is insufficient, test-time strategies may not be able to compensate.

❷ Limited Domain Generalization:

1. **Dynamic Domain Routing:** Use Beam Search Process Reward Model (PRM) to select domain-specific inference paths based on problem types (*e.g.*, detecting the “triangle” keyword and choosing between algebra solvers or geometry theorem provers).
2. **Meta-learning Optimization:** Fine-tune the model on a small number of domain-specific samples via Test-Time Training (TTT) to quickly adapt to new domains (*e.g.*, probability and statistics problems)
3. **Limitations:** Problems with blurred domain boundaries (*e.g.*, math application problems involving common sense reasoning) may fail due to routing errors.

❸ Error Feedback Limitations:

1. **Process Supervision Reinforcement:** Use PRM to validate each step of reasoning in real-time. If an error is detected (*e.g.*, misuse of integration symbols), backtrack and correct the path; combine Self-Consistency by generating multiple inference paths and selecting the one with no contradictions via majority voting.
2. **Limitations:** The reliability of PRM depends on the coverage of error types in the training data. Long-tail errors such as rare symbol confusions may be overlooked.

In summary, combining the flexibility of test-time scaling with the specialization of multimodal tools can help mitigate the core challenges in multimodal mathematical reasoning. However, **it is crucial to balance computational efficiency and accuracy.**