

Solving Inequality Proofs with Large Language Models

Jiayi Sheng^{*} ^β, Luna Lyu^{*} ^α, Jikai Jin^α, Tony Xia^α, Alex Gu^γ, James Zou[†] ^α, Pan Lu[†] ^α

^α Stanford University ^β UC Berkeley ^γ Massachusetts Institute of Technology

Website: <https://ineqmath.github.io/>

 Code  Dataset  Leaderboard

Abstract

Inequality proving, crucial across diverse scientific and mathematical fields, tests advanced reasoning skills such as discovering tight bounds and strategic theorem application. This makes it a distinct, demanding frontier for large language models (LLMs), offering insights beyond general mathematical problem-solving. Progress in this area is hampered by existing datasets that are often scarce, synthetic, or rigidly formal. We address this by proposing an *informal yet verifiable* task formulation, recasting inequality proving into two automatically checkable subtasks: bound estimation and relation prediction. Building on this, we release INEQMATH, an expert-curated dataset of Olympiad-level inequalities, including a test set and training corpus enriched with step-wise solutions and theorem annotations. We also develop a novel LLM-as-judge evaluation framework, combining a *final-answer* judge with four *step-wise* judges designed to detect common reasoning flaws. A systematic evaluation of 29 leading LLMs on INEQMATH reveals a surprising reality: even top models like o1 achieve less than 10% overall accuracy under step-wise scrutiny; this is a drop of up to 65.5% from their accuracy considering only final answer equivalence. This discrepancy exposes fragile deductive chains and a critical gap for current LLMs between merely finding an answer and constructing a rigorous proof. Scaling model size and increasing test-time computation yield limited gains in overall proof correctness. Instead, our findings highlight promising research directions such as theorem-guided reasoning and self-refinement.

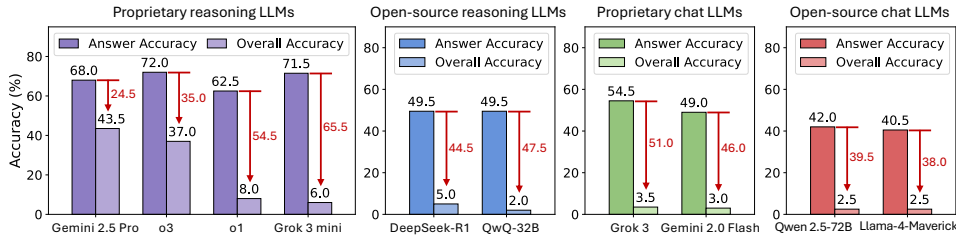


Figure 1: Final-answer accuracy versus overall accuracy for leading LLMs across different categories on the INEQMATH benchmark of Olympiad-level inequality problems. Overall accuracy, measuring both answer correctness and step soundness, is substantially lower than final-answer accuracy for all model types. This highlights a critical gap: while LLMs may find correct final answers to these inequality problems, their reasoning is often unsound. Each model used its optimal maximal tokens.

1 Introduction

Mathematical inequalities are fundamental to diverse fields such as analysis, optimization, and probability theory, with applications spanning scientific modeling, economics, and competitive mathematics. Proving an inequality is a complex endeavor, demanding not just calculation but a sophisticated blend of intuition for discovering tight bounds, strategic insight for selecting and applying classical theorems (e.g., AM-GM, Cauchy-Schwarz), and precise symbolic transformations. These skills are hallmarks of advanced mathematical reasoning, distinguishing inequality proving from general math problem-solving. Automating this process would therefore have broad impact: it could supply automated theorem provers (ATPs) with missing lemmas, accelerate formal verification

^{*} Co-first authors. [†] Co-senior authors. ✉{jamesz, panlu}@stanford.edu

processes, and serve as a demanding testbed for general-purpose reasoners. However, despite impressive advancements in LLMs like DeepSeek-R1 (DeepSeek-AI, 2025) and OpenAI o3 (OpenAI, 2025b), as well as in ATPs themselves (Dong et al., 2024; Hu et al., 2025; Lin et al., 2025a; Poesia et al., 2024; Ye et al., 2024), automating inequality proving remains a challenging frontier.

A major bottleneck in advancing LLM capabilities for inequality proving is the scarcity of suitable benchmarks. Existing resources fall short in several ways: general ATP collections like MiniF2F (Zheng et al., 2022) and ProofNet (Azerbayev et al., 2024) contain few inequalities; synthetic datasets such as INT (Wu et al., 2021) and AIPS (Wei et al., 2024) offer scale but may lack structural diversity due to template-based generation; and curated collections like ChenNEQ (Chen, 2014) are often too small for extensive training. More fundamentally, most existing datasets adopt a *fully formal* representation, where problems and proofs are encoded in systems like Lean (de Moura et al., 2015) or Isabelle (Nipkow et al., 2002). While formal mathematical reasoning offers correctness guarantees and is a vital research direction, LLMs, trained on vast corpora of natural language, often exhibit strong informal reasoning capabilities. This suggests LLMs might solve problems informally even when struggling with the exacting syntax of formal provers. Our work, therefore, aims to explore and benchmark these informal abilities, complementing formal mathematical AI by focusing on a mode of reasoning closer to human intuition, often less structured, stages of mathematical discovery.

To bridge this gap between formal rigor and intuitive problem-solving, we propose an *informal yet verifiable* formulation (§2). Rather than requiring full, machine-checkable proofs within formal systems, we reformulate inequality problems into two concrete, automatically verifiable subtasks: (i) *Bound estimation*—determine the largest (or smallest) constant C that preserves the inequality; and (ii) *Relation prediction*—identify which relation ($>$, $\geq$, $=$, $\leq$, or $<$) holds between two expressions. Both tasks can be presented in natural language and $\text{\LaTeX}$, solved step-by-step by an LLM, and their final answers (a constant or a relation symbol) can be automatically checked. This preserves the creative essence of inequality proving while avoiding the heavy overhead of formal proof assistants.

Building on this formulation, we present INEQMATH (§3), the first large-scale dataset of Olympiad-level inequalities written entirely in informal language. The *test set* comprises 200 original problems, each crafted and reviewed by IMO-level medalists to ensure both originality and difficulty. The *training corpus* includes 1,252 problems sourced from advanced textbooks, automatically rephrased by LLMs into our subtasks and then meticulously reviewed by human experts. A key feature is that each training problem is accompanied by up to four *step-wise solution paths*, providing rich data for training LLMs on fine-grained reasoning. Additionally, 76.8% of the training problems are annotated with 83 *named theorems* across 29 categories relevant to their solutions. As shown in Table 2, INEQMATH surpasses prior resources in scale, diversity, and alignment with human-like, informal problem-solving approaches.

However, producing the correct final answer is insufficient; the reasoning process itself must be sound. To assess this, we introduce an *LLM-as-judge* evaluation framework (§4). This framework comprises a high-precision *final-answer judge* to verify the answer equivalence, complemented by four specialized *step-wise judges* for step soundness. These step-wise judges are designed to detect the frequent reasoning flaws identified in our pilot studies: inappropriate reliance on *toy case* examples, unaddressed *logical gaps*, unjustified *numeric approximations*, and *numeric calculation* errors. Validated on manually labeled solutions, these judges demonstrate high reliability ($F1 > 0.9$ on average) and offer a scalable method to scrutinize the deductive integrity of proofs.

We evaluate 29 leading LLMs ranging from chat models to advanced reasoning LLMs, both open-source and proprietary (§5). As key results highlighted in Figure 1, several key findings emerge. While specialized reasoning LLMs (e.g., o1 (OpenAI, 2024c)) achieve higher *final-answer* accuracy than general-purpose chat models (e.g., GPT-4o (OpenAI, 2024a)), this advantage often collapses under step-wise scrutiny. Once our judges inspect every reasoning step, *overall* accuracy plummets by up to 65.5%. Indeed, even top-performing models like o1 achieve less than 10% overall accuracy (Table 4), exposing fragile deductive chains and a significant gap between finding an answer and constructing a rigorous proof.

Our in-depth study (§5.3) reveals that while larger model sizes correlate with improved *final-answer accuracy*, their impact on *overall accuracy* is limited (e.g., o1 achieves only 8.0% overall accuracy). Similarly, extending test-time computation through longer reasoning chains offers diminishing returns for overall correctness (e.g., o1’s 8.0% overall accuracy remains unchanged when scaling max

completion tokens from 5K to 40K, while o3 (OpenAI, 2025b) saturates around 31%). These findings suggest that current scaling approaches are insufficient for robust deductive reasoning in INEQMATH. Instead, we explore promising improvement strategies, demonstrating potential gains from methods such as theorem-guided reasoning—by providing golden theorems (improving overall accuracy by up to 11% for o3-mini (OpenAI, 2025)) and critic-guided self-refinement (e.g., a 5% absolute increase in overall accuracy for Gemini 2.5 Pro (Google DeepMind, 2025b)).

In summary, our work makes four key contributions: 1) We introduce an *informal* reformulation of inequality proving, decomposing the task into two verifiable subtasks (§2). 2) We release INEQMATH, an expert-curated benchmark of Olympiad-level inequalities and a training corpus enriched with step-wise solutions and theorem annotations (§3). 3) We develop a modular *LLM-as-judge* framework that rigorously evaluates both final answers and proof step soundness (§4). 4) We conduct a systematic empirical study (§5) that exposes a pronounced gap between LLM performance and mathematical rigor, highlighting avenues for future research.

2 Task formalization: an informal perspective

Inequality proof problems require demonstrating that a specified inequality holds under given conditions, such as proving $a + b \geq 2\sqrt{ab}$ for all positive real numbers a and b . Traditionally, these problems are formalized in proof assistants like Lean or Isabelle, represented as a tuple (S_0, I, P) , where S_0 is the initial state, I is the inequality, and P is a set of premises. The proof process, often modeled as a Markov Decision Process, constructs a step-by-step solution verified by the system. However, this formal approach demands expertise in specialized tools, while informal proofs in natural language, though more intuitive, are difficult to verify automatically due to their unstructured nature.

To address these challenges, we propose an *informal* perspective that reformulates inequality proof problems into two *verifiable* subtasks: **bound estimation** and **relation prediction**.

INEQMATH Training Example 1: Bound Problem

Question: Find the maximal constant C such that for all real numbers a, b, c , the inequality holds:

$$\sqrt{a^2 + (1-b)^2} + \sqrt{b^2 + (1-c)^2} + \sqrt{c^2 + (1-a)^2} \geq C$$

Solution: Applying Minkowsky’s Inequality to the left-hand side we have

$$\sqrt{a^2 + (1-b)^2} + \sqrt{b^2 + (1-c)^2} + \sqrt{c^2 + (1-a)^2} \geq \sqrt{(a+b+c)^2 + (3-a-b-c)^2}$$

By denoting $a + b + c = x$, we get

$$\sqrt{(a+b+c)^2 + (3-a-b-c)^2} = \sqrt{2\left(x - \frac{3}{2}\right)^2 + \frac{9}{2}} \geq \sqrt{\frac{9}{2}} = \boxed{\frac{3\sqrt{2}}{2}}.$$

Minkowsky’s Inequality Theorem: For any real number $r \geq 1$ and any positive real numbers $a_1, a_2, \dots, a_n, b_1, b_2, \dots, b_n$

$$\left(\sum_{i=1}^n (a_i + b_i)^r\right)^{\frac{1}{r}} \leq \left(\sum_{i=1}^n a_i^r\right)^{\frac{1}{r}} + \left(\sum_{i=1}^n b_i^r\right)^{\frac{1}{r}}$$

This **bound estimation** task involves finding an optimal constant for a given inequality. For example, in $a + b \geq C\sqrt{ab}$ for $\forall a, b > 0$, the objective is to find the largest C . Formally, a bound estimation problem instance is a triple:

$$\Pi_{\text{bound}} = (f(\mathbf{x}), g(\mathbf{x}), \mathcal{D}), \quad \text{where } \mathcal{D} \subseteq \mathbb{R}^n.$$

Here, $f, g : \mathcal{D} \rightarrow \mathbb{R}$ are two expressions involving variables $\mathbf{x} = (x_1, \dots, x_n)$ within a specified domain $\mathcal{D}$ (e.g., $x_i > 0, \sum x_i = 1$), and $g(\mathbf{x}) > 0, \forall \mathbf{x} \in \mathcal{D}$. The goal is to determine the extremal:

$$C^* = \sup\{C \in \mathbb{R} : f(\mathbf{x}) \geq Cg(\mathbf{x}), \forall \mathbf{x} \in \mathcal{D}\} \quad \text{or} \quad C^* = \inf\{C \in \mathbb{R} : f(\mathbf{x}) \leq Cg(\mathbf{x}), \forall \mathbf{x} \in \mathcal{D}\}.$$

The **relation prediction** task requires determining the correct relationship between two expressions. For instance, given expressions $f(\mathbf{x}) = a + b$ and $g(\mathbf{x}) = 2\sqrt{ab}$, the goal is to identify the relation

INEQMATH Training Example 2: Relation Problem

Question: Let a, b, c be positive real numbers such that $abc = 1$. Consider the following expressions:

$$\frac{b+c}{\sqrt{a}} + \frac{c+a}{\sqrt{b}} + \frac{a+b}{\sqrt{c}} \quad () \quad \sqrt{a} + \sqrt{b} + \sqrt{c} + 3$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Solution: From the AM-GM Inequality, we have

$$\begin{aligned} \frac{b+c}{\sqrt{a}} + \frac{c+a}{\sqrt{b}} + \frac{a+b}{\sqrt{c}} &\geq 2 \left(\sqrt{\frac{bc}{a}} + \sqrt{\frac{ca}{b}} + \sqrt{\frac{ab}{c}} \right) \\ &= \left(\sqrt{\frac{bc}{a}} + \sqrt{\frac{ca}{b}} \right) + \left(\sqrt{\frac{ca}{b}} + \sqrt{\frac{ab}{c}} \right) + \left(\sqrt{\frac{ab}{c}} + \sqrt{\frac{bc}{a}} \right) \\ &\geq 2(\sqrt{a} + \sqrt{b} + \sqrt{c}) \boxed{\geq} \sqrt{a} + \sqrt{b} + \sqrt{c} + 3\sqrt[6]{abc} = \sqrt{a} + \sqrt{b} + \sqrt{c} + 3. \end{aligned}$$

AM-GM Inequality Theorem: If $a_1, a_2, \dots, a_n$ are nonnegative real numbers, then

$$\frac{1}{n} \sum_{i=1}^n a_i \geq (a_1 a_2 \dots a_n)^{\frac{1}{n}}$$

with equality if and only if $a_1 = a_2 = \dots = a_n$. This is a special case of the Power Mean Inequality.

(in this case, $\geq$) that holds for $\forall a, b > 0$. Formally, a relation prediction problem instance is a triple:

$$\Pi_{\text{rel}} = (f(\mathbf{x}), g(\mathbf{x}), \mathcal{D}),$$

where $f(\mathbf{x})$ and $g(\mathbf{x})$ are expressions over variables $\mathbf{x}$ in domain $\mathcal{D} \subseteq \mathbb{R}^n$. The goal is to find the relation between $f(\mathbf{x})$ and $g(\mathbf{x})$ (i.e. $>$, $\geq$, $=$, $\leq$, $<$, or none of the above).

These subtasks are chosen because they frequently appear in mathematical problem-solving, simplify the evaluation process, and crucially, retain the core reasoning challenges inherent in original inequality proof problems. An ideal LLM solution should not only produce the correct final answer but also present a clear, logically sound, and complete derivation. This includes strategic application of theorems, accurate symbolic manipulations and calculations, and justification of all critical steps.

3 INEQMATH: the inequality problem dataset

This section describes the data curation and statistics of INEQMATH, a novel collection of inequality problems designed to support the informal perspective on solving and proving inequalities.

Test data curation. To mitigate contamination from common sources (textbooks, contests, online resources) potentially in LLM training corpora, we commissioned IMO-level medalists to design novel inequality problems. These underwent rigorous review by a separate expert group and were validated only upon unanimous confirmation of solvability, soundness, and ground truth correctness. Problems identified as easier by experts were excluded from the test set (repurposed for development) to ensure a high challenge level. See the developed curation tool in §B.1. We host an online leaderboard¹ with automatic submission and evaluation, providing a reliable and fair community platform.

Training data curation. Training problems were sourced from two advanced textbooks featuring graduate-level and Olympiad-style inequality proof problems. We parsed these textbooks to extract proof problems, their step-wise solutions, and relevant theorems. We developed two LLM-based rephrasers to transform each source problem into two sub-tasks defined in §2: bound estimation and relation prediction. For instance, a source problem like “Prove $a + b \geq 2\sqrt{ab}$ for $\forall a, b \in \mathbb{R}^+$ ” would be rephrased into a bound estimation task (e.g., “Determine the maximal constant C such that $a + b \geq C\sqrt{ab}$ holds for $\forall a, b \in \mathbb{R}^+$ ”) and a relation prediction task (e.g., “Determine the inequality relation in the expression $a + b () 2\sqrt{ab}$ that holds for $\forall a, b \in \mathbb{R}^+$ ”).

¹<https://huggingface.co/spaces/AI4Math/IneqMath-Leaderboard>

Crucially, while rephrased problems are altered from the source proof problem in the format, they preserve the core mathematical reasoning and solution steps—such as applying relevant theorems, determining boundary conditions, and verifying inequalities. An annotation tool (see §B.1) was developed to facilitate human expert review and correction of the LLM-rephrased problems. Extracted theorems were curated, each including its name, a natural language definition, and a list of training problems where it is applicable.

Key statistics. As shown in Table 1, the INEQMATH dataset comprises 200 test problems for benchmarking, 100 development problems with public ground truth, and 1,252 training problems split evenly between bound estimation and relation prediction tasks. Each training problem includes step-wise solutions, with up to four solutions per problem, and 76.8% (962 problems) are annotated with relevant theorems. The dataset features 83 named theorems across 29 categories, with their distribution illustrated in Figure 2. Test problem examples are provided in §B.3.

Statistic	Number	Bnd.	Rel.
Theorem categories	29	-	-
Named theorems	83	-	-
Training problems (for training)	1252	626	626
- With theorem annotations	962	482	480
- With solution annotations	1252	626	626
- Avg. solutions per problem	1.05	1.06	1.05
- Max solutions per problem	4	4	4
Dev problems (for development)	100	50	50
Test problems (for benchmarking)	200	96	104

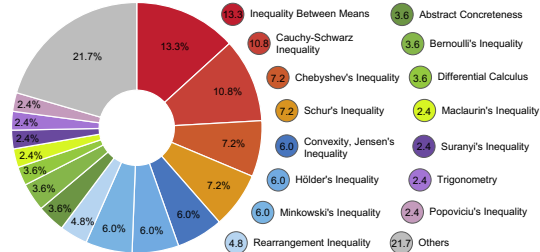


Table 1: Statistics of the INEQMATH dataset.

Figure 2: Distribution of theorem categories.

Comparison to existing datasets. As summarized in Table 2, INEQMATH stands out for: (1) providing expert-curated training and test sets, (2) offering rich annotations with step-wise solutions and 83 grounded theorems, and (3) adopting an informal, accessible format for inequality proving through bound estimation and relation prediction, evaluated via LLM-as-judge. This design bridges the gap between formal proof systems and intuitive mathematical reasoning, making INEQMATH a unique resource for advancing LLM capabilities in problem solving and theorem proving.

Potential contamination statement. To ensure rigorous evaluation, the INEQMATH test set was commissioned from IMO-level medalists to feature novel problems, minimizing prior LLM pre-training exposure. The poor performance across models (§5.2), particularly on overall accuracy (which demands step-wise correctness), strongly suggests the benchmark poses a significant reasoning challenge, regardless of potential familiarity with underlying mathematical concepts. We therefore believe the INEQMATH test set effectively probes novel problem-solving capabilities, and our conclusions on current LLM limitations in rigorous inequality proving remain robust.

Datasets	Data Source		Data Annotation		Problem and Evaluation		
	Training	Test / Dev	#Theorem	Solution	Category	Format	Evaluation
INT (Wu et al., 2021)	Synthesized	Synthesized	35	✓	Proof	Formal	Symbolic DSL
AIPS (Wei et al., 2024)	Synthesized	✗	8	✓	Proof	Formal	Symbolic DSL
MO-INT (Wei et al., 2024)	✗	Data compilation	✗	✗	Proof	Formal	Symbolic DSL
MINIF2F (Zheng et al., 2022)	✗	Autoformalization	✗	✗	Proof	Formal	LEAN
ProofNet (Azerbayev et al., 2024)	✗	Autoformalization	✗	✗	Proof	Formal	LEAN
FormalMATH (Yu et al., 2025)	✗	Autoformalization	✗	✗	Proof	Formal	LEAN
leanWorkbook (Ying et al., 2024)	Autoformalization	Autoformalization	✗	✗	Proof	Formal	LEAN
Proof or Bluff (Petrov et al., 2025)	✗	Data compilation	✗	✗	Proof	Informal	Human judge
CHAMP (Mao et al., 2024)	✗	Autoformalization	✗	✗	Open	Informal	Human judge
Putnam Axiom (Gulati et al., 2024)	✗	Data compilation	✗	✗	Open	Informal	Answer checking
LiveMathBench (Liu et al., 2024)	✗	Data compilation	✗	✗	Open	Informal	Answer checking
INEQMATH (Ours)	Expert annotated	Expert annotated	83	✓	MC, Open	Informal	LLM-as-judge

Table 2: Comparison of datasets for inequality and theorem proving. INEQMATH provides expert-annotated training and test/dev sets, featuring high-quality named theorems and step-wise solutions for model development. Unlike prior datasets using synthesis or autoformalization, INEQMATH presents problems in informal language across multiple-choice (MC) and open-ended (Open) formats, and employs LLM-as-judge for evaluation.

4 Fine-grained informal judges for inequality solving

The test split of the INEQMATH dataset serves as our benchmark, comprising 200 Olympiad-level inequality problems that challenge both humans and current LLMs. Traditional evaluation methods fall

short in this setting: expert annotation is accurate but prohibitively labor-intensive, while automated techniques such as string matching or value equivalence fail to capture step-by-step correctness—an essential aspect of inequality problem solving. To address this, we propose a fine-grained *LLM-as-judge* framework, consisting of a *final-answer judge* for verifying the predicted answer (§4.1) and four specialized *step-wise judges* targeting common reasoning flaws (§4.2). A solution is deemed correct *overall* only if it passes all five judges. As shown in Table 3, these judges achieve strong alignment with human annotations ($F1 = 0.93$), providing a scalable yet reliable alternative to manual evaluation.

LLM-as-Judge	Judge type	Accuracy	Precision	Recall	F1 score
Final Answer Judge	Answer checking	1.00	1.00	1.00	1.00
Toy Case Judge	Step soundness	0.91	0.86	0.97	0.91
Logical Gap Judge	Step soundness	0.96	0.95	0.98	0.96
Numerical Approximation Judge	Step soundness	0.96	0.95	0.98	0.96
Numerical Computation Judge	Step soundness	0.71	0.68	0.98	0.80
Average	-	0.91	0.89	0.98	0.93

Table 3: Performance metrics of *LLM-as-judge* framework on development set.

4.1 Final answer judge

LLM-generated solutions to INEQMATH problems typically involve multiple reasoning steps followed by a concluding answer statement. However, the final answer may vary in phrasing or numeric format, especially for bound estimation problems. For example, $C = \frac{1}{\sqrt{2}}$ and $C = \frac{\sqrt{2}}{2}$ are mathematically equivalent but differ in form. Recent work (Lu et al., 2024) evaluates LLM outputs via format normalization and exact string matching, without accounting for mathematical equivalence. To address this, we propose a two-stage *Final Answer Judge*: it first identifies the concluding sentence with the predicted answer, then performs robust equivalence checking, even when the form differs from the reference. Prompt details and examples are in §C.2.

4.2 Four step-wised judges

Toy case judge. Inequality problems in INEQMATH often require reasoning over continuous domains (e.g., all $a, b, c > 0$), where specific numerical examples alone are insufficient for a valid proof. LLM frequently generalize incorrectly from such examples—e.g., claiming an inequality holds universally because it holds for $a = 1, b = 2, c = 3$. Prior work (Gao et al., 2025) flags these under a broad “logical flaw” category, lacking granularity for targeted analysis. Our *Toy Case Judge* addresses this by detecting unjustified generalization from toy examples. It prompts an LLM to flag conclusions based solely on specific instances without broader justification. See §C.3 for prompts and examples.

Logical gap judge. INEQMATH inequality problems often involve multi-step derivations (e.g., algebraic manipulation, constrained optimization, functional transformations) needing explicit justification. LLMs, however, often skip key reasoning steps or assert conclusions without support (e.g., stating an optimal bound without derivation). Existing step-level evaluations (Xia et al., 2024) assess validity and redundancy but lack granularity for such logical omissions. Our *Logical Gap Judge* addresses this by flagging missing transitions, unjustified claims, and vague derivations, especially in steps involving inequality transformations or bound estimation (see §C.4 for details).

Numerical approximation judge. Inequality problems in INEQMATH often demand exact symbolic reasoning, where the use of numeric approximations—e.g., replacing $\sqrt{2}$ with 1.414—can compromise mathematical rigor. However, many LLM-generated solutions resort to such approximations during intermediate steps, leading to inaccurate or non-generalizable conclusions. To address this, we introduce a *Numerical Approximation Judge* that flags inappropriate use of numeric approximations—specifically when they affect derivations or final answers. Approximations used solely for intuition or side remarks are allowed. See §C.5 for prompt details and examples.

Numerical computation judge. Many INEQMATH problems require explicit numerical computations after variable assignment (e.g., evaluating $\frac{27}{2}$ or summing rational terms). While symbolic reasoning is vital, arithmetic accuracy is equally crucial for overall correctness. Prior work (e.g., EIC-Math (Li et al., 2024a)) categorizes broad error types but often overlooks subtle miscalculations in multi-step derivations. Our *Numerical Computation Judge* addresses this by verifying arithmetic steps once variables are instantiated. It prompts an LLM to extract numerical expressions, convert them into Python code, and evaluate using floating-point arithmetic within a small tolerance. This enables

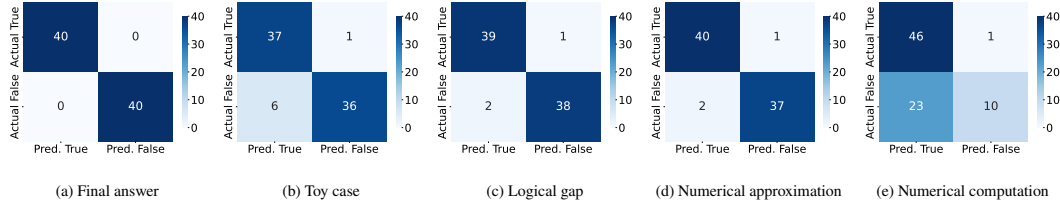


Figure 4: Confusion matrices for five judges, which exhibit strong agreement with human labels.

high-precision checking of both intermediate and final results. To further improve precision and mitigate floating-point issues, we encourage using symbolic mathematics packages such as SymPy, particularly for handling fractions and decimal numbers. Additional details are provided in §C.6.

4.3 Effectiveness verification of judges

A holistic LLM judge baseline. To motivate our specialized judging system, we first evaluate a heuristic LLM-as-judge baseline. This prompts a single, general-purpose LLM to holistically assess INEQMATH solution correctness, based on both final answer accuracy and step-wise soundness across the four reasoning categories in §4.2. As shown in the confusion matrix (Figure 3) using 80 human-annotated development examples, this naive approach exhibits poor agreement with human labels, underscoring its unreliability for rigorous evaluation in this domain.

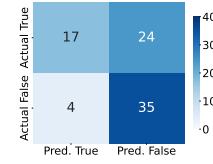


Figure 3: Confusion matrix for the judge baseline.

Performance of our fine-grained judges. In contrast, our proposed suite of five specialized judges exhibits strong alignment with human evaluations. Figure 4 presents the confusion matrices for each judge on the same development set. The final answer judge (using GPT-4o-mini) achieves near-perfect agreement, while the four step-wise judges (chosen for a balance of performance and cost as detailed in §C.7) also demonstrate high fidelity. This confirms that decomposing the complex evaluation task into targeted sub-problems allows LLMs to serve as reliable evaluators.

5 Experiments in INEQMATH

5.1 Experimental setups

We conduct a systematic evaluation of 29 leading LLMs on the inequality problems in the INEQMATH test set. The evaluated models span two categories: general-purpose chat models (both open-source and proprietary) and specialized reasoning LLMs designed for complex, multi-step problem-solving. All models are prompted in a zero-shot setting with the problem statement and the instruction: “*Please solve the problem with clear, rigorous, and logically sound steps*” to encourage detailed reasoning. Model responses are assessed using our LLM-as-judge framework (§4). We report three key metrics:

- *Answer Acc*: Measures the predicted answer correctness, verified by the final-answer judge (§4.1).
- *Step Acc*: Aggregates the correctness of individual reasoning steps as determined by our four specialized step-wise judges (§4.2), which target common flaws.
- *Overall Acc*: The primary metric, which deems a solution correct only if it achieves both a correct final answer and flawless step-wise reasoning (i.e., passes all five judges).

A response is thus considered fully correct (*Overall Acc*) only if it produces the right final answer through logically valid steps, passing scrutiny from all judges. Additional setup details are in §D.1.

5.2 Main evaluation results

Table 4 presents the performance of the evaluated LLMs on INEQMATH. Our analysis reveals several critical insights into current LLM capabilities for inequality proving:

1) Reasoning LLMs achieve higher final-answer accuracy. Models like o1 (62.5% *Answer Acc*) and Grok 3 mini (71.5% *Answer Acc*) significantly outperform their general-purpose chat counterparts (e.g., GPT-4o at 37.5%, Grok 3 at 54.5%) in identifying the correct final answer. This suggests specialized architectures or training for reasoning improve their search ability to find answers.

2) Step-wise scrutiny reveals a dramatic performance drop. The advantage in *Answer Acc* often masks underlying reasoning flaws. *Overall Acc* plummets when steps are evaluated. For instance,

Models	Overall Acc (↑)			Answer Acc (↑)			Step Acc (↑)											
							No Toy Case			No Logic. Gap			No Approx. Error			No Comp. Error		
	All	Bnd.	Rel.	All	Bnd.	Rel.	All	Bnd.	Rel.	All	Bnd.	Rel.	All	Bnd.	Rel.	All	Bnd.	Rel.
Heuristic Methods																		
Random Guess	-	-	-	8.5	0.0	16.7	-	-	-	-	-	-	-	-	-	-	-	-
Frequent Guess	-	-	-	18.0	9.4	26.0	-	-	-	-	-	-	-	-	-	-	-	-
Open-source Chat LLMs																		
Qwen2.5-Coder-32B (Hui et al., 2024)	1.5 ^{139.0}	1.0 ^{30.0}	1.9 ^{28.9}	40.5	<u>51.0</u>	30.8	36.0	27.1	44.2	3.0	2.1	3.8	90.5	96.9	84.6	88.5	89.6	87.5
Llama-4-Scout (Meta Platforms, Inc., 2025b)	1.5 ^{132.0}	2.1 ^{144.8}	1.0 ^{30.2}	33.5	46.9	21.2	30.5	15.6	44.2	3.5	<u>4.2</u>	2.9	<u>93.0</u>	94.8	<u>91.3</u>	92.5	92.7	92.3
Qwen2.5-72B (Qwen Team, 2024a)	2.5 ^{139.5}	3.1 ^{147.9}	1.9 ^{31.8}	<u>42.0</u>	<u>51.0</u>	33.7	<u>54.5</u>	<u>53.1</u>	<u>55.8</u>	<u>5.0</u>	<u>4.2</u>	<u>5.8</u>	91.0	94.8	87.5	<u>95.0</u>	<u>94.8</u>	95.2
Llama-4-Maverick (Meta Platforms, Inc., 2025a)	2.5 ^{138.0}	2.1 ^{143.7}	2.9 ^{32.7}	40.5	45.8	<u>35.6</u>	42.5	28.1	<u>55.8</u>	4.0	<u>4.2</u>	3.8	89.0	91.7	86.5	<u>95.0</u>	92.7	<u>97.1</u>
Qwen2.5-7B (Qwen Team, 2024b)	3.0 ^{132.0}	2.1 ^{138.5}	<u>3.8</u> ^{28.0}	35.0	40.6	29.8	44.5	32.3	<u>55.8</u>	4.5	3.1	<u>5.8</u>	92.5	<u>96.9</u>	88.5	93.0	92.7	93.3
Proprietary Chat LLMs																		
Gemini 2.0 Flash-Lite (Google DeepMind, 2025d)	1.5 ^{131.5}	2.1 ^{141.7}	1.0 ^{22.1}	33.0	43.8	23.1	11.5	11.5	11.5	3.5	3.1	3.8	73.0	77.1	69.2	90.5	87.5	93.3
GPT-4o mini (OpenAI, 2024b)	2.0 ^{137.5}	1.0 ^{141.7}	2.9 ^{33.6}	39.5	42.7	36.5	29.0	11.5	<u>45.2</u>	2.5	2.1	2.9	90.0	91.7	88.5	93.0	92.7	93.3
GPT-4.1 (OpenAI, 2025a)	2.5 ^{138.0}	0.0 ^{131.3}	<u>4.8</u> ^{44.2}	40.5	31.3	49.0	16.0	12.0	19.0	10.0	8.3	11.5	59.5	66.7	52.9	93.5	92.7	<u>94.2</u>
GPT-4o (OpenAI, 2024a)	3.0 ^{134.5}	2.1 ^{138.5}	3.8 ^{30.8}	37.5	40.6	34.6	<u>32.0</u>	<u>21.9</u>	43.0	3.5	3.1	3.8	<u>92.5</u>	<u>93.8</u>	<u>91.4</u>	94.0	93.8	<u>94.2</u>
Gemini 2.0 Flash (Google DeepMind, 2025c)	3.0 ^{146.0}	3.1 ^{156.3}	2.9 ^{36.5}	49.0	59.4	39.4	15.5	13.5	17.3	13.5	7.3	19.2	55.5	60.4	51.0	<u>94.5</u>	<u>94.8</u>	<u>94.2</u>
Grok 3 (xAI, 2025a)	3.5 ^{151.0}	<u>4.2</u> ^{162.5}	2.9 ^{40.4}	<u>54.5</u>	<u>66.7</u>	43.3	17.0	13.7	20.2	<u>16.0</u>	<u>11.6</u>	<u>20.2</u>	36.0	42.1	30.8	93.0	<u>96.8</u>	90.4
Open-source Reasoning LLMs																		
QwQ-32B (Alibaba Qwen Team, 2025)	2.0 ^{147.5}	2.1 ^{152.1}	1.9 ^{143.3}	49.5	54.2	45.2	26.0	25.0	26.9	29.5	20.1	37.5	21.0	20.8	21.2	87.0	82.3	<u>91.3</u>
Deepseek-R1 (Llama-70B) (DeepSeek-AI, 2025a)	3.5 ^{150.0}	5.2 ^{153.1}	1.9 ^{147.1}	<u>53.5</u>	<u>58.3</u>	49.0	23.0	24.0	22.1	26.0	20.9	30.8	<u>35.5</u>	<u>38.5</u>	32.7	87.0	89.6	84.6
Deepseek-R1 (Qwen-14B) (DeepSeek-AI, 2025b)	5.0 ^{135.5}	<u>6.3</u> ^{136.4}	3.8 ^{134.7}	40.5	42.7	38.5	21.0	18.8	23.1	21.0	19.8	22.1	<u>35.5</u>	<u>38.5</u>	32.7	85.0	91.7	78.8
Deepseek-R1 (DeepSeek-AI, 2025)	5.0 ^{144.5}	4.2 ^{163.5}	5.8 ^{136.9}	49.5	67.7	32.7	57.0	53.1	60.9	17.5	6.3	27.9	81.0	95.8	67.3	95.0	<u>99.0</u>	91.3
Qwen3-235B-A22B (Qwen Team, 2025)	<u>6.0</u> ^{135.0}	3.1 ^{132.3}	<u>8.7</u> ^{137.5}	41.0	35.4	46.2	<u>35.0</u>	<u>30.2</u>	<u>39.4</u>	<u>36.0</u>	<u>26.0</u>	<u>45.2</u>	31.0	28.1	<u>33.7</u>	<u>92.5</u>	<u>93.8</u>	<u>91.3</u>
Proprietary Reasoning LLMs																		
Claude 3.7 Sonnet (Anthropic, 2025)	2.0 ^{140.0}	2.1 ^{144.8}	1.9 ^{135.6}	42.0	46.9	37.5	49.0	36.5	60.6	4.0	3.1	4.8	93.5	95.8	91.3	93.0	90.6	95.2
Gemini 2.5 Flash (Google DeepMind, 2025a)	4.5 ¹¹⁰	3.1 ¹¹¹	5.8 ^{10.9}	5.5	4.2	6.7	88.0	84.4	91.3	13.5	7.3	19.2	<u>100.0</u>	<u>100.0</u>	<u>100.0</u>	<u>100.0</u>	<u>100.0</u>	<u>100.0</u>
Grok 3 mini (xAI, 2025b)	6.0 ^{165.5}	4.2 ^{168.7}	7.7 ^{162.5}	<u>71.5</u>	<u>72.9</u>	<u>70.2</u>	24.0	16.7	30.8	19.5	11.5	26.9	<u>53.5</u>	63.5	44.2	91.0	94.8	87.5
Gemini 2.5 Pro (Google DeepMind, 2025b)	6.0 ¹¹⁰	7.3 ¹¹⁰	4.8 ¹¹⁰	7.0	8.3	5.8	88.5	83.3	93.3	19.0	12.5	25.0	<u>100.0</u>	<u>100.0</u>	<u>100.0</u>	<u>99.5</u>	<u>100.0</u>	99.0
o1 (OpenAI, 2024c)	8.0 ^{154.5}	7.3 ^{155.2}	8.7 ^{153.8}	62.5	62.5	62.5	34.5	37.5	31.7	17.5	12.5	22.1	86.5	99.0	75.0	<u>99.5</u>	<u>100.0</u>	99.0
o3-mini (OpenAI, 2025)	9.5 ^{153.0}	7.3 ^{162.5}	11.5 ^{143.3}	62.5	69.8	55.8	37.0	34.4	39.4	22.0	17.7	26.0	77.5	92.7	63.5	95.0	96.9	93.3
o4-mini (OpenAI, 2025b)	15.5 ^{149.5}	14.6 ^{148.9}	16.3 ^{150.0}	65.0	63.5	66.3	62.0	58.3	65.4	26.0	25.0	26.9	86.5	90.6	82.7	93.0	92.7	93.3
o3 (OpenAI, 2025b)	<u>21.0</u> ^{116.0}	<u>18.8</u> ^{114.4}	<u>23.1</u> ¹⁰²	37.0	30.2	43.3	<u>93.5</u>	<u>91.7</u>	<u>95.2</u>	<u>39.5</u>	<u>28.1</u>	<u>50.0</u>	91.5	99.0	84.6	97.0	96.9	97.1
Average Accuracy (↑)	5.0 ^{138.0}	4.5 ^{142.9}	5.5 ^{135.5}	43.0	47.4	39.0	40.3	34.8	45.5	15.0	11.0	18.7	73.1	77.9	68.8	93.2	93.7	92.8
Average Error Rate (↓)	95.0 ^{138.0}	95.5 ^{142.9}	94.5 ^{135.5}	57.0	52.6	61.0	59.7	65.2	54.5	85.0	89.0	81.3	26.9	22.1	31.2	6.8	6.3	7.2

Table 4: Evaluation performance of chat and reasoning LLMs on the INEQMATH benchmark (the test set). *Bnd.* denotes bound problems and *Rel.* denotes relation ones. We report: (1) *Overall Acc.*, which reflects the correctness of both the final answer and intermediate steps; (2) *Answer Acc.*, which measures final answer correctness alone; and (3) *Step Acc.*, which evaluates the accuracy of intermediate steps across four error categories—*Toy Case*, *Logical Gap*, *Numerical Approximation*, and *Numerical Computation*. **Blue superscripts** ↓ indicate accuracy drop (*Overall Acc* - *Answer Acc*) from step-wise errors. Underlining denotes best result within each model category; **boldface** highlights best overall performance. Default max token limit for reasoning LLMs is 10K.

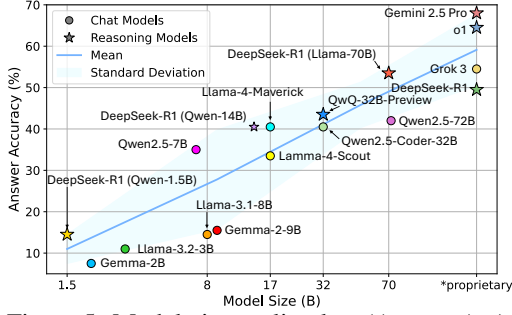
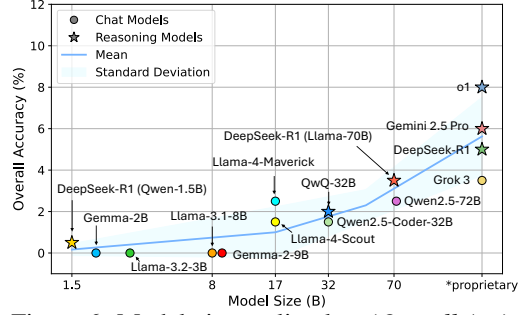
Grok 3 mini’s accuracy drops by 65.5% (from 71.5% *Answer Acc* to 6.0% *Overall Acc*), and o3-mini by 53.0%. This stark discrepancy underscores the fragility of LLM-generated deductive chains.

3) Robust proof construction remains a major challenge. Even top models like o1 achieve low *Overall Acc* (8.0%). Many large models, despite moderate *Answer Acc*, also score poorly (e.g., Grok 3 at 3.5% *Overall Acc*). This indicates a fundamental gap between finding a plausible answer and constructing a mathematically rigorous, step-by-step derivation.

5.3 In-depth study

Failure solution analysis. As shown in Table 4, where we report average error rates for overall accuracy, final-answer accuracy, and step-wise accuracy across four categories, the most common step-wise errors in LLM-generated solutions are logical gaps (85.0% average failure rate across models) and improper generalization from toy cases (59.7%). Less frequent, but still significant, are errors from numerical approximations (26.9%) and miscalculations (6.8%). A detailed examination of incorrect solutions (see examples in §D.2.1-§D.2.4) highlights these prevalent error patterns, which often undermine proofs even when LLMs produce the correct final answer. Beyond these step-wise errors, LLMs also struggle to derive correct final answers for complex problems (§D.2.5), indicating deeper challenges in theorem application and symbolic manipulation.

Scaling law in model size. Figure 5 shows how *final-answer* accuracy (which evaluates only the correctness of the final predicted answer) scales with model size for LLMs. As model size increases, we observe a steady improvement in answer accuracy, reflecting an empirical scaling law that larger models are better at inferring correct bounds and inequality relationships. However, this trend does not hold well when considering *overall accuracy*—which requires both a correct answer and valid intermediate reasoning steps—as shown in Figure 6. In this latter case, the scaling curve flattens, indicating that increased model size alone is insufficient to eliminate step-by-step reasoning errors.

Figure 5: Model-size scaling law (*Answer Acc*).Figure 6: Model-size scaling law (*Overall Acc*).

Scaling law in test-time computation. Extended test-time computation, allowing longer reasoning chains, is a common strategy for complex problem-solving (DeepSeek-AI, 2025). We investigated its impact on *overall accuracy* in INEQMATH by varying the maximum completion tokens for reasoning LLMs. Figure 7 shows that while models like Gemini 2.5 Pro and o3 initially improve with more tokens, performance gains saturate (e.g., beyond 20K tokens). This indicates that merely increasing computational budget offers diminishing returns for achieving rigorous, step-wise correct proofs, highlighting the need for more than just longer thought processes.

5.4 Exploring improvement strategies

Retrieving relevant theorems as hints. To assess theorem-based hints, we provide models with the top- k most frequent theorems from our INEQMATH training corpus when solving a 40-problem test subset. As shown in Figure 8, providing one or two such theorems decreases *overall accuracy* for weaker models (e.g., Grok 3 mini, o3-mini, o4-mini), likely due to misapplication or distraction by potentially irrelevant information. Conversely, stronger models like Gemini 2.5 Pro benefit from these hints, suggesting advanced reasoning is crucial to effectively use such guidance. These results underscore the potential of theorem-guided reasoning but also highlight the critical need for more sophisticated theorem retrieval mechanisms (e.g., RAG (Lewis et al., 2020; Gupta et al., 2024)) to reliably enhance LLM performance in inequality proving. Detailed experiments are available in §D.4.

Self-improvement via critic as feedback. Allowing an LLM to critique and revise its own reasoning has been shown to improve performance on complex tasks (Yuksekgonul et al., 2025; Tian et al., 2024). To explore whether this holds for inequality proving, we drew 40 random test problems from INEQMATH and ran one round of self-critique. As Figure 9 shows, self-critique consistently improves performance—e.g., Gemini 2.5 Pro’s overall accuracy rises from 43% to 48%. This trend underscores self-critique as a promising method to enhance logical rigor and solution quality of LLMs in inequality reasoning. More details are in §D.5.

6 Conclusion

In summary, we introduce an informal yet verifiable task formulation for inequality proving, decomposing it into bound estimation and relation prediction. Building on this, we release INEQMATH, an expert-curated benchmark of Olympiad-level inequalities with a training corpus featuring step-wise solutions and theorem annotations. Our novel LLM-as-judge evaluation framework, comprising a final-answer judge and four step-wise judges, enables a rigorous assessment. Our comprehensive evaluation of diverse leading LLMs reveals a critical gap: while LLMs may achieve high final-answer accuracy, this often plummets by up to 65.5% under step-wise scrutiny, with top models like o1 achieving less than 10% overall accuracy. This discrepancy exposes fragile deductive chains for current LLMs in constructing rigorous proofs. We further find that scaling model size or increasing test-time computation yields limited gains in overall proof correctness. Instead, our findings highlight promising research directions such as theorem-guided reasoning and self-refinement.

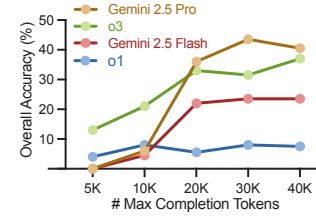


Figure 7: Scaling law in test-time computation for reasoning LLMs.

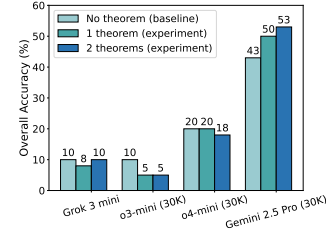


Figure 8: Model performance with retrieved theorems as hints.

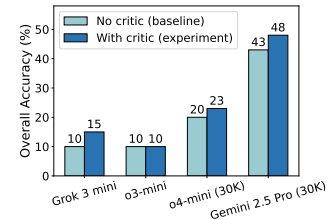


Figure 9: Model performance via self-critic as feedback.

Acknowledgments

This work is partially supported by the Hoffman-Yee Research Grants program at Stanford HAI. We thank Yu (Bryan) Zhou for early discussion and feedback.

References

- Meta AI. Llama 3.1 8b. <https://huggingface.co/meta-llama/Llama-3.1-8B>, 2024a. Accessed: 2025-05-15.
- Meta AI. Llama 3.2 3b. <https://huggingface.co/meta-llama/Llama-3.2-3B>, 2024b. Accessed: 2025-05-15.
- Alibaba Qwen Team. QwQ-32B. <https://huggingface.co/Qwen/QwQ-32B>, 2025. Hugging Face model card.
- Anthropic. Claude 3.7 Sonnet. <https://www.anthropic.com/claude/sonnet>, 2025. Anthropic model card.
- Zhangir Azerbayev, Bartosz Piotrowski, Hailey Schoelkopf, Edward William Ayers, and Dragomir Radev. Proofnet: Autoformalizing and formally proving undergraduate-level mathematics, 2024. URL <https://openreview.net/forum?id=Zix86UbMGh>.
- Evan Chen. A brief introduction to olympiad inequalities, 2014. URL <https://web.evanchen.cc/handouts/Ineq/en.pdf>. Accessed: 2025-03-19.
- Wenhu Chen, Ming Yin, Max Ku, Pan Lu, Yixin Wan, Xueguang Ma, Jianyu Xu, Xinyi Wang, and Tony Xia. TheoremQA: A theorem-driven question answering dataset. In Houda Bouamor, Juan Pino, and Kalika Bali (eds.), *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 7889–7901, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.489. URL <https://aclanthology.org/2023.emnlp-main.489/>.
- Leonardo Mendonça de Moura, Soonho Kong, Jeremy Avigad, Floris van Doorn, and Jakob von Raumer. The lean theorem prover (system description). In *CADE*, 2015. URL <https://api.semanticscholar.org/CorpusID:232990>.
- DeepSeek-AI. DeepSeek-R1-Distill-Llama-70B. <https://huggingface.co/deepseek-ai/DeepSeek-R1-Distill-Llama-70B>, 2025a. Hugging Face model card.
- DeepSeek-AI. DeepSeek-R1-Distill-Qwen-14B. <https://huggingface.co/deepseek-ai/DeepSeek-R1-Distill-Qwen-14B>, 2025b. Hugging Face model card.
- DeepSeek-AI. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, 2025. URL <https://arxiv.org/abs/2501.12948>.
- Kefan Dong and Tengyu Ma. Beyond limited data: Self-play llm theorem provers with iterative conjecturing and proving. *arXiv preprint arXiv:2502.00212*, 2025.
- Kefan Dong, Arvind V. Mahankali, and Tengyu Ma. Formal theorem proving by rewarding LLMs to decompose proofs hierarchically. In *The 4th Workshop on Mathematical Reasoning and AI at NeurIPS’24*, 2024. URL <https://openreview.net/forum?id=D83tiHiNfF>.
- Bofei Gao, Feifan Song, Zhe Yang, Zefan Cai, Yibo Miao, Qingxiu Dong, Lei Li, Chenghao Ma, Liang Chen, Runxin Xu, Zhengyang Tang, Benyou Wang, Daoguang Zan, Shanghaoran Quan, Ge Zhang, Lei Sha, Yichang Zhang, Xuancheng Ren, Tianyu Liu, and Baobao Chang. Omni-MATH: A universal olympiad level mathematic benchmark for large language models. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=yagPf0KAlN>.
- Fabian Gloeckle, Jannis Limperg, Gabriel Synnaeve, and Amaury Hayat. ABEL: Sample efficient online reinforcement learning for neural theorem proving. In *The 4th Workshop on Mathematical Reasoning and AI at NeurIPS’24*, 2024. URL <https://openreview.net/forum?id=kk3mSjVCUO>.

- Google DeepMind. Gemini 2.5 Flash. <https://cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-5-flash>, 2025a. Vertex AI model card.
- Google DeepMind. Gemini 2.5 Pro. <https://deepmind.google/technologies/gemini/pro/>, 2025b. Google DeepMind model card.
- Google DeepMind. Gemini 2.0 Flash. <https://cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-0-flash>, 2025c. Vertex AI model card.
- Google DeepMind. Gemini 2.0 Flash-Lite. <https://cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-0-flash-lite>, 2025d. Vertex AI model card.
- Aryan Gulati, Brando Miranda, Eric Chen, Emily Xia, Kai Fronsdal, Bruno de Moraes Dumont, and Sanmi Koyejo. Putnam-AXIOM: A functional and static benchmark for measuring higher level mathematical reasoning. In *The 4th Workshop on Mathematical Reasoning and AI at NeurIPS'24*, 2024. URL <https://openreview.net/forum?id=YXnwlZe0yf>.
- Shailja Gupta, Rajesh Ranjan, and Surya Narayan Singh. A comprehensive survey of retrieval-augmented generation (rag): Evolution, current landscape and future directions. *arXiv preprint arXiv:2410.12837*, 2024.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring mathematical problem solving with the MATH dataset. In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*, 2021. URL <https://openreview.net/forum?id=7Bywt2mQsCe>.
- Jiewen Hu, Thomas Zhu, and Sean Welleck. miniCTX: Neural theorem proving with (long-)contexts. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=KigaAqEFHW>.
- Binyuan Hui, Jian Yang, Zeyu Cui, Jiayi Yang, Dayiheng Liu, Lei Zhang, Tianyu Liu, Jiajun Zhang, Bowen Yu, Keming Lu, et al. Qwen2. 5-coder technical report. *arXiv preprint arXiv:2409.12186*, 2024.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. Retrieval-augmented generation for knowledge-intensive nlp tasks. In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 9459–9474. Curran Associates, Inc., 2020. URL https://proceedings.neurips.cc/paper_files/paper/2020/file/6b493230205f780e1bc26945df7481e5-Paper.pdf.
- Wenda Li, Lei Yu, Yuhuai Wu, and Lawrence C Paulson. Isarstep: a benchmark for high-level mathematical reasoning. In *International Conference on Learning Representations*, 2021.
- Xiaoyuan Li, Wenjie Wang, Moxin Li, Junrong Guo, Yang Zhang, and Fuli Feng. Evaluating mathematical reasoning of large language models: A focus on error identification and correction. *arXiv preprint arXiv:2406.00755*, 2024a.
- Yang Li, Dong Du, Linfeng Song, Chen Li, Weikang Wang, Tao Yang, and Haitao Mi. Hunyuanprover: A scalable data synthesis framework and guided tree search for automated theorem proving. *arXiv preprint arXiv:2412.20735*, 2024b.
- Zenan Li, Zhaoyu Li, Wen Tang, Xian Zhang, Yuan Yao, Xujie Si, Fan Yang, Kaiyu Yang, and Xiaoxing Ma. Proving olympiad inequalities by synergizing LLMs and symbolic reasoning. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=FiyS0ecSm0>.
- Zhenwen Liang, Linfeng Song, Yang Li, Tao Yang, Feng Zhang, Haitao Mi, and Dong Yu. Mps-prover: Advancing stepwise theorem proving by multi-perspective search and data curation. *arXiv preprint arXiv:2505.10962*, 2025.
- Haohan Lin, Zhiqing Sun, Sean Welleck, and Yiming Yang. Lean-STar: Learning to interleave thinking and proving. In *The Thirteenth International Conference on Learning Representations*, 2025a. URL <https://openreview.net/forum?id=SOWZ59UyNc>.

- Yong Lin, Shange Tang, Bohan Lyu, Jiayun Wu, Hongzhou Lin, Kaiyu Yang, Jia Li, Mengzhou Xia, Danqi Chen, Sanjeev Arora, et al. Goedel-prover: A frontier model for open-source automated theorem proving. *arXiv preprint arXiv:2502.07640*, 2025b.
- Junnan Liu, Hongwei Liu, Linchen Xiao, Ziyi Wang, Kuikun Liu, Songyang Gao, Wenwei Zhang, Songyang Zhang, and Kai Chen. Are your llms capable of stable reasoning? *arXiv preprint arXiv:2412.13147*, 2024.
- Pan Lu, Hritik Bansal, Tony Xia, Jiacheng Liu, Chunyuan Li, Hannaneh Hajishirzi, Hao Cheng, Kai-Wei Chang, Michel Galley, and Jianfeng Gao. Mathvista: Evaluating mathematical reasoning of foundation models in visual contexts. In *The Twelfth International Conference on Learning Representations*, 2024.
- Yujun Mao, Yoon Kim, and Yilun Zhou. CHAMP: A competition-level dataset for fine-grained analyses of LLMs’ mathematical reasoning capabilities. In *Findings of the Association for Computational Linguistics: ACL 2024*, pp. 13256–13274, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.findings-acl.785. URL <https://aclanthology.org/2024.findings-acl.785/>.
- Meta Platforms, Inc. Llama-4-Maverick-17B-128E. <https://huggingface.co/meta-llama/Llama-4-Maverick-17B-128E>, 2025a. Hugging Face model card.
- Meta Platforms, Inc. Llama-4-Scout-17B-16E. <https://huggingface.co/meta-llama/Llama-4-Scout-17B-16E>, 2025b. Hugging Face model card; accessed 2025-05-12.
- Tobias Nipkow, Markus Wenzel, and Lawrence C. Paulson. *Isabelle/HOL: A Proof Assistant for Higher-Order Logic*. Springer, 2002. ISBN 978-3-540-43376-7. doi: <https://doi.org/10.1007/3-540-45949-9>.
- OpenAI. GPT-4o. <https://openai.com/index/hello-gpt-4o/>, 2024a. OpenAI blog post.
- OpenAI. GPT-4o mini. <https://openai.com/index/gpt-4o-mini-advancing-cost-efficient-intelligence/>, 2024b. OpenAI blog post.
- OpenAI. OpenAI o1. <https://openai.com/o1/>, 2024c. OpenAI official announcement.
- OpenAI. Openai o3-mini system card, January 2025. URL <https://cdn.openai.com/o3-mini-system-card-feb10.pdf>. Accessed: 2025-03-19.
- OpenAI. GPT-4.1. <https://openai.com/index/gpt-4-1/>, 2025a. OpenAI model announcement.
- OpenAI. OpenAI o4-mini. <https://openai.com/index/introducing-o3-and-o4-mini/>, 2025b. OpenAI official announcement.
- Ivo Petrov, Jasper Dekoninck, Lyuben Baltadzhiev, Maria Drencheva, Kristian Minchev, Mislav Balunović, Nikola Jovanović, and Martin Vechev. Proof or bluff? evaluating llms on 2025 usa math olympiad. *arXiv preprint arXiv:2503.21934*, 2025.
- Gabriel Poesia, David Broman, Nick Haber, and Noah Goodman. Learning formal mathematics from intrinsic motivation. *Advances in Neural Information Processing Systems*, 37:43032–43057, 2024.
- Qwen Team. Qwen2.5-72B. <https://huggingface.co/Qwen/Qwen2.5-72B>, 2024a. Hugging Face model card.
- Qwen Team. Qwen2.5-7B. <https://huggingface.co/Qwen/Qwen2.5-7B>, 2024b. Hugging Face model card.
- Qwen Team. Qwen3-235B-A22B. <https://huggingface.co/Qwen/Qwen3-235B-A22B>, 2025. Hugging Face model card.
- Z. Z. Ren, Zhihong Shao, Junxiao Song, Huajian Xin, Haocheng Wang, Wanjia Zhao, Liyue Zhang, Zhe Fu, Qihao Zhu, Dejian Yang, Z. F. Wu, Zhibin Gou, Shirong Ma, Hongxuan Tang, Yuxuan Liu, Wenjun Gao, Daya Guo, and Chong Ruan. Deepseek-prover-v2: Advancing formal mathematical reasoning via reinforcement learning for subgoal decomposition, April 2025.

Gemma Team, Thomas Mesnard, Cassidy Hardin, Robert Dadashi, Surya Bhupatiraju, Shreya Pathak, Laurent Sifre, Morgane Rivière, Mihir Sanjay Kale, Juliette Love, et al. Gemma: Open models based on gemini research and technology. *arXiv preprint arXiv:2403.08295*, 2024.

Ye Tian, Baolin Peng, Linfeng Song, Lifeng Jin, Dian Yu, Lei Han, Haitao Mi, and Dong Yu. Toward self-improvement of LLMs via imagination, searching, and criticizing. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=tPdJ2qHk0B>.

George Tsoukalas, Jasper Lee, John Jennings, Jimmy Xin, Michelle Ding, Michael Jennings, Amitayush Thakur, and Swarat Chaudhuri. Putnambench: Evaluating neural theorem-provers on the putnam mathematical competition. *arXiv preprint arXiv:2407.11214*, 2024.

Nguyen Duy Tung. 567 nice and hard inequalities, April 2012. URL <https://phamtuankhai.wordpress.com/wp-content/uploads/2012/04/567-bat-dang-thuc-hay.pdf>. Accessed: 2025-03-19.

Haiming Wang, Huajian Xin, Chuanyang Zheng, Zhengying Liu, Qingxing Cao, Yinya Huang, Jing Xiong, Han Shi, Enze Xie, Jian Yin, Zhenguo Li, and Xiaodan Liang. LEGO-prover: Neural theorem proving with growing libraries. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=3f5PALef5B>.

Haiming Wang, Mert Unsal, Xiaohan Lin, Mantas Baksys, Junqi Liu, Marco Dos Santos, Flood Sung, Marina Vinyes, Zhenzhe Ying, Zekai Zhu, Jianqiao Lu, Hugues de Saxe’e, Bolton Bailey, Chendong Song, Chenjun Xiao, Dehao Zhang, Ebony Zhang, Frederick Pu, Han Zhu, Jiawei Liu, Jonas Bayer, Julien Michel, Longhui Yu, Léo Dreyfus-Schmidt, Lewis Tunstall, Luigi Pagani, Moreira Machado, Pauline Bourigault, Ran Wang, Stanislas Polu, Thibaut Barroyer, Wen-Ding Li, Yazhe Niu, Yann Fleureau, Yangyang Hu, Zhouliang Yu, Zihan Wang, Zhilin Yang, Zhengying Liu, and Jia Li. Kimina-prover preview: Towards large formal reasoning models with reinforcement learning. *arXiv preprint arXiv:2504.11354*, 2025.

Chenrui Wei, Mengzhou Sun, and Wei Wang. Proving olympiad algebraic inequalities without human demonstrations. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=8kFctyli9H>.

Yuhuai Wu, Albert Jiang, Jimmy Ba, and Roger Baker Grosse. {INT}: An inequality benchmark for evaluating generalization in theorem proving. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=O6LPudownQm>.

Zijian Wu, Suozhi Huang, Zhejian Zhou, Huaiyuan Ying, Jiayu Wang, Dahua Lin, and Kai Chen. Internlm2. 5-stepprover: Advancing automated theorem proving via expert iteration on large-scale lean problems. *arXiv preprint arXiv:2410.15700*, 2024.

xAI. Grok 3. <https://x.ai/news/grok-3>, 2025a. xAI official announcement.

xAI. Grok 3 Mini. <https://x.ai/news/grok-3>, 2025b. xAI official announcement.

Shijie Xia, Xuefeng Li, Yixin Liu, Tongshuang Wu, and Pengfei Liu. Evaluating mathematical reasoning beyond accuracy. *CoRR*, abs/2404.05692, 2024. URL <https://doi.org/10.48550/arXiv.2404.05692>.

Huajian Xin, Daya Guo, Zhihong Shao, Z.Z. Ren, Qihao Zhu, Bo Liu, Chong Ruan, Wenda Li, and Xiaodan Liang. Advancing theorem proving in LLMs through large-scale synthetic data. In *The 4th Workshop on Mathematical Reasoning and AI at NeurIPS’24*, 2024. URL <https://openreview.net/forum?id=TPtXLikhny>.

Ran Xin, Chenguang Xi, Jie Yang, Feng Chen, Hang Wu, Xia Xiao, Yifan Sun, Shen Zheng, and Kai Shen. Bfs-prover: Scalable best-first tree search for llm-based automatic theorem proving. *arXiv preprint arXiv:2502.03438*, 2025.

Kaiyu Yang and Jia Deng. Learning to prove theorems via interacting with proof assistants. In *International Conference on Machine Learning*, pp. 6984–6994. PMLR, 2019.

- Kaiyu Yang, Aidan M Swope, Alex Gu, Rahul Chalamala, Peiyang Song, Shixing Yu, Saad Godil, Ryan Prenger, and Anima Anandkumar. Leandojo: Theorem proving with retrieval-augmented language models. In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2023. URL <https://openreview.net/forum?id=g70X2s0Jtn>.
- Ziyu Ye, Jiacheng Chen, Jonathan Light, Yifei Wang, Jiankai Sun, Guohao Li, Mac Schwager, Philip Torr, Yuxin Chen, Kaiyu Yang, Yisong Yue, and Ziniu Hu. Reasoning in reasoning: A hierarchical framework for (better and faster) neural theorem proving. In *The 4th Workshop on Mathematical Reasoning and AI at NeurIPS'24*, 2024. URL <https://openreview.net/forum?id=H5hePMXKht>.
- Huaiyuan Ying, Zijian Wu, Yihan Geng, Jiayu Wang, Dahua Lin, and Kai Chen. Lean workbook: A large-scale lean problem set formalized from natural language math problems. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=Vcw3vzjHDb>.
- Zhouliang Yu, Ruotian Peng, Keyi Ding, Yizhe Li, Zhongyuan Peng, Minghao Liu, Yifan Zhang, Zheng Yuan, Huajian Xin, Wenhao Huang, Yandong Wen, Ge Zhang, and Weiyang Liu. Formalmath: Benchmarking formal mathematical reasoning of large language models, 2025. URL <https://arxiv.org/abs/2505.02735>.
- Mert Yuksekgonul, Federico Bianchi, Jack Boen, et al. Optimizing generative ai by backpropagating language model feedback. *Nature*, 639:609–616, 2025. doi: 10.1038/s41586-025-08661-4. URL <https://doi.org/10.1038/s41586-025-08661-4>. Published 19 March 2025.
- Ziyin Zhang, Jiahao Xu, Zhiwei He, Tian Liang, Qiuzhi Liu, Yansi Li, Linfeng Song, Zhengwen Liang, Zhuosheng Zhang, Rui Wang, et al. Deeptheorem: Advancing llm reasoning for theorem proving through natural language and reinforcement learning. *arXiv preprint arXiv:2505.23754*, 2025.
- Haoyu Zhao, Yihan Geng, Shange Tang, Yong Lin, Bohan Lyu, Hongzhou Lin, Chi Jin, and Sanjeev Arora. Ineq-Comp: Benchmarking human-intuitive compositional reasoning in automated theorem proving on inequalities. *arXiv preprint arXiv:2505.12680*, 2025a.
- Zhiyu Zhao, Yongcheng Zeng, Ning Yang, and Guoqing Liu. Enhancing mathematical reasoning in language models through focused differentiation training, 2025b. URL <https://openreview.net/forum?id=vuuYbA1vB2>.
- Kunhao Zheng, Jesse Michael Han, and Stanislas Polu. minif2f: a cross-system benchmark for formal olympiad-level mathematics. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=9ZPegFuFTFv>.

Supplementary Materials for Solving Inequality Proofs with Large Language Models

Appendix Contents

A	Related Work	16
B	Dataset Curation Details	17
B.1	Data Annotation Tool	17
B.2	Prompts for Rephrasing Problems	18
B.3	INEQMATH Examples	20
C	Fine-grained Informal Judge Details	22
C.1	Qualitative analysis of judge disagreements.	22
C.2	Final Answer Judge	22
C.3	Toy Case Judge	24
C.4	Logical Gap Judge	25
C.5	Numerical Approximation Judge	27
C.6	Numerical Computation Judge	28
C.7	Development Performance of Judges	30
C.8	Judge Failure Examples	31
D	Experimental Details for Inequality Solving	35
D.1	Experimental Setups	35
D.2	Model Failure Solution Examples	37
D.3	Taking Annotated Theorems as Hints	43
D.4	Retrieval as Augmentation	46
D.5	Self-improvement via Critic as Feedback	51
E	Limitations	53
F	Broader Impacts	53

A Related Work

Datasets for inequality and theorem proving. One of major bottlenecks in advancing LLM capabilities for inequality proving is the scarcity of suitable datasets. Existing resources fall short in several ways: general ATP collections like MiniF2F (Zheng et al., 2022) and ProofNet (Azerbayev et al., 2024) contain few inequalities; synthetic datasets such as INT (Wu et al., 2021) and AIPS (Wei et al., 2024) offer scale but often lack structural diversity due to their template-based generation; and curated collections like ChenNEQ (Chen, 2014) are often too small for extensive training. More fundamentally, most existing datasets (Zhao et al., 2025a; Tung, 2012; Yang & Deng, 2019; Li et al., 2021; Tsoukalas et al., 2024; Hu et al., 2025) adopt a *fully formal* representation, where problems and proofs are encoded in systems like Lean (de Moura et al., 2015) or Isabelle (Nipkow et al., 2002). While formal mathematical reasoning offers correctness guarantees and is a vital research direction, LLMs, trained on vast corpora of natural language, often exhibit strong informal reasoning capabilities. Therefore, our INEQMATH adopts an *informal* perspective, reformulating inequality proof problems into two verifiable subtasks—bound estimation and relation prediction. These problems within INEQMATH were crafted and reviewed by IMO-level medalist experts. Other informal reasoning datasets (Petrov et al., 2025; Mao et al., 2024; Gulati et al., 2024; Liu et al., 2024) typically lack annotated solutions, theorem references, or corresponding training data. To address these gaps, INEQMATH introduces 1,252 inequality problems for training, each annotated with theorems (from a set of 83 named theorems spanning 29 categories) relevant to its solution, which comprises up to four steps.

Methods for inequality and theorem proving. Proving inequalities is a complex endeavor, demanding intuition for discovering tight bounds, strategic insight for selecting and applying classical theorems, and precise symbolic transformations. Traditional automated theorem provers (ATPs) primarily operate within formal systems like Lean (de Moura et al., 2015) or Isabelle (Nipkow et al., 2002), requiring problems and proofs to be encoded in specialized languages. Inspired by the mathematical reasoning capabilities of LLMs (Zhao et al., 2025b), a significant body of recent work has focused on integrating LLMs with these formal ATPs. These approaches often model theorem proving as a Markov Decision Process (MDP), training LLMs to select appropriate tactics and premises to construct proofs verifiable by the formal system (Dong et al., 2024; Gloeckle et al., 2024; Hu et al., 2025; Lin et al., 2025a; Poesia et al., 2024; Ye et al., 2024; Xin et al., 2024; Li et al., 2025; Liang et al., 2025; Zhang et al., 2025; Dong & Ma, 2025; Wang et al., 2025; Ren et al., 2025; Wang et al., 2024). For instance, systems like Goedel-Prover (Lin et al., 2025b) leverage large Lean corpora to train models for tactic prediction, enabling end-to-end formal proof generation. Other methods incorporate tree-search techniques to navigate the vast search space of premises within these formal frameworks (Wu et al., 2024; Li et al., 2024b; Xin et al., 2025).

However, LLMs are fundamentally trained on vast corpora of natural language, suggesting their inherent strengths may lie in informal reasoning, which is often closer to human intuition and the preliminary stages of mathematical discovery. This highlights a gap and an opportunity for methods that leverage these informal capabilities. Our work diverges from the purely formal paradigm by proposing an *informal yet verifiable* approach to inequality proving, aiming to benchmark and enhance LLM performance in a setting that better aligns with their human-like problem-solving, while also exploring potential improvement strategies like theorem-guided reasoning and self-refinement.

LLM-as-judge for math problem solving. Reliable evaluation of mathematical problem-solving necessitates assessing not only final answer correctness but also the logical soundness of each reasoning step, a significant challenge for automated systems. Traditional methods are often inadequate: expert annotation is labor-intensive and unscalable for large-scale evaluation (Petrov et al., 2025; Mao et al., 2024), while automated techniques such as string matching or value equivalence overlook crucial step-by-step proof correctness (Hendrycks et al., 2021; Gulati et al., 2024; Liu et al., 2024; Lu et al., 2024). While LLMs show promise as evaluators (LLM-as-judge), their capacity for detailed, step-wise mathematical judgment is still developing. Existing step-level LLM judges (Xia et al., 2024; Gao et al., 2025), for instance, may assess general step validity but often lack the granularity to identify specific, nuanced reasoning flaws. Similarly, frameworks like EIC-Math (Li et al., 2024a) provide broad error categories but can miss subtle yet critical issues in multi-step derivations, such as minor miscalculations or unstated assumptions. To address these limitations and rigorously assess informal mathematical proofs like inequality solving, our *LLM-as-judge* framework combines a high-precision *final-answer* judge with four *step-wise* judges targeting common errors: toy case overgeneralization, logical gaps, unjustified numeric approximations, and numeric calculation mistakes.

B Dataset Curation Details

B.1 Data Annotation Tool

Problem Checker & Editor
PREVIOUS
train_bound/example.json
NEXT

Problem

Let $a, b, c \in \mathbb{R}^+$. Find the smallest constant C such that the following inequality holds for all a, b, c :

$$C \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$$

Let $a, b, c \in \mathbb{R}^+$. Find the smallest constant C such that the following inequality holds for all a, b, c :

$$C \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$$

Answer

$C = 3$

$C = 3$

Source Proof Problem

Let $a, b, c \in \mathbb{R}^+$. Prove the inequality $3 \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$.

Let $a, b, c \in \mathbb{R}^+$. Prove the inequality $3 \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$.

Source Proof Solution

Let $a \geq b \geq c$. Then $a^3 \geq b^3 \geq c^3$ and $a^5 \geq b^5 \geq c^5$. Due to Chebishev's inequality we have

$$3 \left(a^3 + b^3 + c^3 \right) \left(a^5 + b^5 + c^5 \right) \geq 3 \left(a^8 + b^8 + c^8 \right)$$

Let $a \geq b \geq c$. Then $a^3 \geq b^3 \geq c^3$ and $a^5 \geq b^5 \geq c^5$. Due to Chebishev's inequality we have

$$\left(a^3 + b^3 + c^3 \right) \left(a^5 + b^5 + c^5 \right) \leq 3 \left(a^8 + b^8 + c^8 \right)$$

Comments

☐ Bad Problem
☐ Unsure
☐ Other Issue

Add any additional comments here...

SAVE
SAVE AND NEXT

Figure 10: The interface of our developed tool for checking and editing the bound problems.

Problem Checker & Editor
PREVIOUS
train_relation/example.json
NEXT

Problem

Let $a, b, c \in \mathbb{R}^+$. Consider the following inequality:

$$3 \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Let $a, b, c \in \mathbb{R}^+$. Consider the following inequality:

$$3 \left(a^3 + b^3 + c^3 \right) \quad () \quad \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Answer

(B) $\geq$

(B) $\geq$

Source Proof Problem

Let $a, b, c \in \mathbb{R}^+$. Prove the inequality $3 \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$.

Let $a, b, c \in \mathbb{R}^+$. Prove the inequality $3 \left(a^3 + b^3 + c^3 \right) \geq \left(a^2 + b^2 + c^2 \right) \left(a^5 + b^5 + c^5 \right)$.

Source Proof Solution

Let $a \geq b \geq c$. Then $a^3 \geq b^3 \geq c^3$ and $a^5 \geq b^5 \geq c^5$. Due to Chebishev's inequality we have

$$3 \left(a^3 + b^3 + c^3 \right) \left(a^5 + b^5 + c^5 \right) \geq 3 \left(a^8 + b^8 + c^8 \right)$$

Let $a \geq b \geq c$. Then $a^3 \geq b^3 \geq c^3$ and $a^5 \geq b^5 \geq c^5$. Due to Chebishev's inequality we have

$$\left(a^3 + b^3 + c^3 \right) \left(a^5 + b^5 + c^5 \right) \leq 3 \left(a^8 + b^8 + c^8 \right)$$

Comments

☐ Bad Problem
☐ Unsure
☐ Other Issue

Add any additional comments here...

SAVE
SAVE AND NEXT

Figure 11: The interface of our developed tool for checking and editing the relation problems.

B.2 Prompts for Rephrasing Problems

Prompt for Rephrasing Proofs to Bound Problems

Task: Transform the given inequality problem into a bound prediction problem by introducing a constant C and determining its optimal value.

Instructions:

1. Analyze the original problem, focusing on its structure and potential for transformation.
2. Introduce a constant C by either replacing an existing constant or creating a new relationship between expressions.
3. Determine whether to find the minimal or maximal value of C that satisfies the inequality for all relevant variables.
4. Consider factors such as homogeneity, existing constraints, and the domain of variables (e.g., positive reals, all reals).
5. Ensure the rephrased problem maintains the mathematical essence and constraints of the original.

Output format:

Provide your response in the following structure:

<Analysis>: Concise explanation of key features and transformation approach.

<Conclusion>: YES or NO, followed by a brief summary of the transformation.

<Rephrased Problem>: Transformed problem statement, focusing on finding the optimal C .

<Answer>: $C = \text{< value >}$.

Key considerations:

1. For double inequalities, attempt to rephrase as a single bound prediction problem when possible.
2. In homogeneous inequalities, focus on the ratios between variables rather than their absolute values.
3. Incorporate any existing constraints into the rephrased version of the problem.
4. Clearly specify the domain of the variables in the rephrased problem statement.
5. Ensure that the rephrased problem is logically equivalent to the original.

Example:

Original problem: Let $a, b, c \in \mathbb{R}^+$. Prove the inequality

$$\frac{abc}{(1+a)(a+b)(b+c)(c+16)} \leq \frac{1}{81}$$

<Analysis>: To turn this into a bound prediction problem, we can focus on the following steps:

1. The left side is a rational expression that is always positive for $a, b, c \in \mathbb{R}^+$.
2. The right side is a fixed constant $\frac{1}{81}$.
3. We replace the constant $\frac{1}{81}$ with a variable C and ask: What is the smallest C such that the inequality holds for all positive a, b, c ?
4. This approach allows us to determine the tightest possible upper bound for the left-hand expression.
5. If we find the smallest C that works, we prove the original inequality and show it's the best possible.

<Conclusion>: YES, the inequality can be rephrased as a bound prediction problem. By replacing the constant $\frac{1}{81}$ with a variable C , we can determine the tightest upper bound for the given rational expression, effectively proving the original inequality and demonstrating its optimality.

<Rephrased problem>:

Coll Determine the minimal constant C such that the following inequality holds for all a, b, c :

$$\frac{abc}{(1+a)(a+b)(b+c)(c+16)} \leq C.$$

<Answer>: $C = \frac{1}{81}$.

Now, please rewrite the following problem:

Original problem: {problem}

Prompt for Rephrasing Proofs to Relation Problems

Task: Transform the given inequality proof problem into a relation prediction problem.

Instructions:

1. Analyze the original problem, identifying key components such as variables, domains, conditions, and the main inequality.
2. Rephrase the problem by maintaining the original expressions and replacing the relation symbol with a blank to be filled.
3. Preserve any additional conditions or constraints from the original problem in your rephrased version.
4. Change the task from “Prove” to “Determine the correct inequality relation to fill in the blank.”
5. Provide a set of options for the relation, always including $\leq$, $\geq$, $=$, $<$, $>$, and “None of the above”.
6. Determine the correct answer based on your modification and analysis.

Output format:

Provide your response in the following structure:

<Analysis>: Detailed step-by-step analysis of the original problem and your approach to rephrasing it.

<Conclusion>: YES or NO, followed by a brief explanation of whether and how the problem can be effectively rephrased.

<Rephrased Problem>:

Transformed problem statement.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

<Answer>: Option letter.

Key considerations:

1. Maintain the original mathematical expressions and any given conditions as much as possible.
2. Ensure the rephrased problem captures the essence and complexity of the original problem.
3. For problems with multiple inequalities, focus on one main inequality for the relation prediction task.
4. When dealing with complex fractions or expressions, keep them intact to maintain the problem’s difficulty level.
5. If the relation depends on specific values of the variables or cannot be definitively determined, consider using “None of the above” as the correct answer.

Example:

Original problem: Let $a, b, c \in \mathbb{R}^+$. Prove the inequality

$$\frac{abc}{(1+a)(a+b)(b+c)(c+16)} \leq \frac{1}{81}$$

<Analysis>: To rephrase it to a relation prediction problem, we can focus on the following steps:

1. The original problem is a proof task for an inequality involving positive real numbers a , b , and c .
2. The left side of the inequality is a complex fraction $\frac{abc}{(1+a)(a+b)(b+c)(c+16)}$.
3. The right side is a constant fraction $\frac{1}{81}$.
4. The original inequality uses the “less than or equal to” ($\leq$) relation, which needs to hold for all positive real values of a , b , and c .
5. We can transform the proof task into determining the correct relation between the left and right sides of the inequality.

<Conclusion>: YES, the inequality can be effectively rephrased as a relation prediction problem.

<Rephrased problem>:

Let $a, b, c \in \mathbb{R}^+$. Consider the following inequality:

$$\frac{abc}{(1+a)(a+b)(b+c)(c+16)} \quad (\quad) \quad \frac{1}{81}.$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

<Answer>: A

Now, please rewrite the following problem:

Original problem: {problem}

B.3 INEQMATH Examples**INEQMATH Testing Example 1: Bound Problem**

Question: Let $x, y, z > 0$ such that $x + y + z = 1$. Determine the minimal constant C such that the following inequality holds for all x, y, z :

$$xy(y + 4z) + yz(z + 4x) + zx(x + 4y) \leq C.$$

INEQMATH Testing Example 2: Bound Problem

Question: Let $a_1, a_2, \dots, a_n$ be real numbers and S be a non-empty subset of $\{1, 2, \dots, n\}$. Find the largest constant C such that the following inequality holds for all $a_1, a_2, \dots, a_n$ and S :

$$2C \left(\sum_{i \in S} a_i \right)^2 \leq \sum_{1 \leq i \leq j \leq n} (a_i + \dots + a_j)^2.$$

INEQMATH Testing Example 3: Bound Problem

Question: Let $a_1, a_2, \dots, a_n > 0$ such that $a_1 + a_2 + \dots + a_n < 1$. Determine the minimal constant C such that the following inequality holds for all $a_1, a_2, \dots, a_n$:

$$\frac{a_1 \cdot a_2 \dots a_n (1 - a_1 - a_2 - \dots - a_n)}{(a_1 + a_2 + \dots + a_n) (1 - a_1) (1 - a_2) \dots (1 - a_n)} \leq C \frac{3}{n^{n-1}}.$$

INEQMATH Testing Example 4: Relation Problem

Question: Let $a, b, c, x, y, z \in \mathbb{R}$ be real numbers such that $a + b + c = 1$ and $x^2 + y^2 + z^2 = 1$. Consider the following expression:

$$a(x + b) + b(y + c) + c(z + a) \quad (\quad) \quad 1.$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

INEQMATH Testing Example 5: Relation Problem

Question: In the plane of the acute-angled triangle $\triangle ABC$, let L be a line such that u, v, w are the lengths of the perpendiculars from A, B, C respectively to L . Consider the following inequality:

$$u^2 \tan A + v^2 \tan B + w^2 \tan C \quad (\quad) \quad \Delta.$$

where Δ is the area of the triangle. Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

INEQMATH Testing Example 6: Relation Problem

Question: Let a, b, c be the sides of any triangle. Consider the following inequality:

$$3 \left(\sum_{cyc} ab(1 + 2 \cos(c)) \right) \quad (\quad) \quad 2 \left(\sum_{cyc} \sqrt{(c^2 + ab(1 + 2 \cos(c))) (b^2 + ac(1 + \cos(b)))} \right).$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

C Fine-grained Informal Judge Details

C.1 Qualitative analysis of judge disagreements.

Despite the strong aggregate performance (overall F1 = 0.93, Table 3), LLM-as-judge evaluations are not perfect. Acknowledging the skepticism surrounding LLM-based evaluation, we conducted a qualitative analysis of failure cases where our judges’ assessments diverged from human annotations. Detailed examples are provided in §C.8. These instances underscore that while highly effective, our LLM judges can still struggle with the deep, nuanced understanding that characterizes expert human mathematical reasoning.

C.2 Final Answer Judge

Prompt for Final Answer Judge: Answer Extraction for Bound problems

You are an expert at extracting numbers from answer sentences. Below are examples of sentences and the corresponding numbers:

Example 1: answer is $C = 2$.

Answer: $C = 2$

Example 2: answer is $C = \frac{1}{\sqrt{2}}$.

Answer: $C = \frac{1}{\sqrt{2}}$

Example 3: answer is $C = 2$.

Answer: $C = 2$

Now, extract the number from the following sentence: {answer_sentence}.

Make sure to return the answer in a format as “C=<extracted_answer>”, where <extracted_answer> is the extracted number or expression.

Prompt for Final Answer Judge: Answer Extraction for Relation Problems

You are an expert at extracting option letters (A, B, C, D, E, F) from answer sentences.

The options are given below:

A: (A) $\leq$

B: (B) $\geq$

C: (C) $=$

D: (D) $<$

E: (E) $>$

F: (F) None of the above

Below are examples of sentences and the corresponding option letters:

Example 1: answer is (B) $\geq$.

Answer: B

Example 2: answer is (E) $>$.

Answer: E

Example 3: answer is: $\leq$.

Answer: A

Now, extract the option letter from the following sentence: {answer_sentence}.

Make sure to return the option letter only, without any other characters.

Prompt for Final Answer Judge: Answer Equivalence Verification

You are an expert at verifying mathematical expression equivalence. Analyze if two expressions are exactly equivalent by following these strict rules:

Required Analysis Steps:

1. Check if both expressions are valid mathematical forms.
2. If either expression is not mathematical (e.g., text or undefined), return False.
3. For numerical expressions:
 - Direct equality (e.g., $2 = 2$) $\rightarrow$ True.
 - Different representations of same value (e.g., $1/2 = 0.5$, $\sqrt{4} = 2$) $\rightarrow$ True.
 - Decimal approximations vs exact values (e.g., $2\pi \neq 6.28318$) $\rightarrow$ False.
4. For algebraic expressions:
 - Must have clear, valid transformation path between forms.
 - If transformation requires multiple non-obvious steps $\rightarrow$ False.
 - Verify equivalence through algebraic proof when possible.
 - For complex expressions, use techniques like squaring or substitution to verify.

Equivalence Criteria:

- Must have exactly same deterministic value.
- Must be provably equivalent through valid mathematical operations.
- Different notations of same exact value are equivalent.
- Decimal approximations are NOT equivalent to exact expressions.
- No rounding or approximations allowed.
- If equivalence cannot be conclusively proven $\rightarrow$ False.

Example pairs and their analysis:

Ground truth: $C = 2$

Prediction: $C = 2$

Analysis: The expressions are identical in both form and value, representing the same integer 2.

Equivalent: True

Ground truth: $C = 1.5$

Prediction: $C = \frac{3}{2}$

Analysis: The decimal 1.5 and fraction $\frac{3}{2}$ are different representations of the same number ($1.5 = \frac{3}{2}$).

Equivalent: True

Ground truth: $C = 2\pi$

Prediction: $C = 6.28318530718$

Analysis: While 6.28318530718 is a decimal approximation of 2π , they are not symbolically equivalent expressions.

Equivalent: False

Ground truth: $C = \sqrt{\frac{1}{6}}$

Prediction: $C = \frac{1}{\sqrt{6}}$

Analysis: These are equivalent through the property $\sqrt{\frac{a}{b}} = \frac{\sqrt{a}}{\sqrt{b}}$ when $a, b > 0$.

Equivalent: True

Ground truth: $C = \sqrt{\frac{3}{2}}$

Prediction: $C = \frac{3}{2\sqrt{2}}$

Analysis: The expressions differ as proven when squared: $(\sqrt{\frac{3}{2}})^2 = \frac{3}{2} \neq \frac{9}{8} = (\frac{3}{2\sqrt{2}})^2$.

Equivalent: False

Now analyze these expressions:

Ground truth: {ground_truth}

Prediction: {prediction}

C.3 Toy Case Judge

Prompt for Toy Case Judge

Task: Evaluate the logical rigor of a solution to an inequality problem, focusing specifically on whether the direction of the inequality was justified using toy cases or special value substitution.

Instructions:

1. Carefully read the reasoning process used to solve the inequality.
2. Identify whether the direction of the inequality was determined by testing special values, trying toy cases, or relying on extreme-case analysis, rather than providing a general proof valid over the entire domain.
3. If the model uses a toy case (e.g., setting a variable to 0, 1, or choosing symmetric/equal values) or considers a variable tending to 0 or ∞ (extreme-case reasoning) to conclude the inequality direction, this should be flagged as logically unsound unless it is later supported by a rigorous or general argument.
4. Substituting special values for the purpose of verifying equality or testing sharpness is acceptable and should not be flagged.
5. If a toy case is used to show that the inequality does not hold (i.e., the two sides are incomparable), this is acceptable and should not be flagged.
6. Trying toy cases or substituting special values for the purpose of exploring or analyzing the problem—without using them to directly conclude the inequality direction—is acceptable and should not be flagged.
7. The goal is to confirm whether the final conclusion is justified for *all variables in the given domain* by using sound and formal reasoning.

Output Format:

<Analysis>: Brief explanation of whether toy cases, special values, or extreme-case reasoning were used to justify inequality direction, and whether this undermines the rigor of the proof.

<Flagged Reasoning Step (if applicable)>: Quote or summarize the step(s) where toy cases or extreme cases were used to justify the inequality.

<Answer>: True or False. It indicates whether the solution's reasoning is logically valid (True) or contains flaws due to unjustified use of toy cases, special value substitution, or extreme-case reasoning (False).

Key Considerations:

1. Check which steps are used to justify the final general conclusion.
2. Toy case examples (e.g., setting , or choosing values like) should not be used to conclude general inequality directions.
3. These techniques are allowed for illustration or equality testing, but not as a replacement for a full justification.
4. Extreme-case analysis (e.g., letting or) used as a substitute for proof to justify the final general conclusion should be flagged.
5. Do not provide improvement suggestions—simply judge whether the logic used is valid or flawed.
6. Only reasoning steps that fail to generalize and is used to justify the final general conclusion should be flagged.

Examples of inputs and outputs:

{examples}

Now analyze the following problem and solution:

Original Problem: {query}

Solution: {response}

Judge Example 1: Toy Case**Response:**

Let's test the specific values.

Case 1: If $a = b = c$, then:

$$\left(\frac{a}{b+c}\right)^2 + \frac{b}{c+a} + \left(\frac{c}{a+b}\right)^2 = \frac{1}{4} + \frac{1}{2} + \frac{1}{4} = 1.$$

Case 2: If $a = 0$ and $b = c = 1$, then:

$$\left(\frac{a}{b+c}\right)^2 + \frac{b}{c+a} + \left(\frac{c}{a+b}\right)^2 = 0 + 1 + 1 = 2.$$

From the test cases, we can conclude that the minimum value of the expression was 1. Therefore, the expression is always at least 1. The answer is $\boxed{C = 1}$.

Judge Output:

<Analysis>: The minimum of the expression is derived from 2 special cases(i.e. $a = b = c$ and $a = 0, b = c = 1$). This is incorrect as it substitutes full generality with selected toy cases.

<Flagged Reasoning Step (if applicable)>: From the test cases, we can conclude that the minimum value of the expression was 1. Therefore, the expression is always at least 1.

<Answer>: $\boxed{False}$

C.4 Logical Gap Judge**Prompt for Logical Gap Judge**

Task: Evaluate the logical rigor of a proposed solution to a mathematical inequality problem. Focus on whether the reasoning includes non-trivial claims made without justification, logical leaps, or unsupported assertions, while allowing for valid optimization-based, algebraic, or analytic analysis when properly demonstrated. This prompt does not evaluate whether the direction of an inequality was justified using toy cases, special values, or asymptotic behavior. That aspect is handled separately.

Instructions:

- Carefully read the entire reasoning process used to solve the inequality.
- Identify whether the solution includes:
 - Any non-obvious (non-trivial) claims or transformations stated without justification or explanation.
 - Any logical gaps or skipped steps that lead to intermediate or final conclusions.
- All significant transformations—especially involving inequalities, bounds, or extremal behavior—must be supported by: Algebraic manipulation, Well-known identities or theorems, Valid analytical tools (e.g., convexity, derivatives, limits) or step-by-step numeric or symbolic reasoning.
- Optimization methods (e.g., Lagrange multipliers, derivative-based analysis) are valid only if the analysis is explicitly shown:
 - If a solution invokes optimization or analytical techniques, it must demonstrate key steps, derivative conditions, or critical point verification.
 - Statements such as “solving the constrained optimization problem confirms...” without any derivation or argument are considered unjustified.
 - You do not need to assess whether toy cases, special values, or extreme behavior were used to infer the inequality direction. That responsibility lies outside the scope of this Judge.
- Simple algebra or widely known transformations (e.g., AM-GM, factoring identities) may be used without full derivation.
- The goal is to assess whether each important conclusion within the reasoning—not just the final answer—is logically supported and rigorously justified.

Prompt for Logical Gap Judge (Continued)**Output Format:**

<Analysis>: Step-by-step explanation of whether the reasoning is logically sound. Highlight any unjustified claims or skipped steps, unless they are supported by valid asymptotic, numeric, or analytic reasoning.

<Flagged Reasoning Step (if applicable)>: Quote or summarize the specific step(s) where an unjustified claim or logical leap occurred.

<Answer>: True or False. True if the reasoning is valid; False if it contains unjustified steps or unsupported claims.

Key Considerations:

1. Focus on whether each major step (not just the final answer) is logically justified.
2. Non-trivial algebraic identities or inequality steps must be explained unless they are well-known.
3. Minor simplifications and standard techniques are acceptable without proof.
4. Do not flag steps involving toy cases, extreme values, or special substitutions used to infer inequality direction—those are out of scope here.
5. Claims like “a numerical check shows” must include specific values, results, or graphs to be valid.
6. Optimization-based arguments (e.g., Lagrange multipliers, critical point methods) must include demonstrated steps or analytic structure. If only the method is named but not applied, the reasoning should be flagged.
7. Do not flag steps used solely for equality verification, sharpness testing, or illustration.
8. Do not provide improvement suggestions—simply judge whether the logic is valid or flawed as presented.

Examples of inputs and outputs:

{examples}

Now analyze the following problem and solution:

Original Problem: {query}

Solution: {response}

Judge Example 2: Logical Gap**Response:**

...

A numerical check confirms that the minimum occurs when $x = 1$, and the expression equals 4. So the largest constant C such that the inequality holds for all positive a, b, c is:

$$C = 4$$

Judge Output:

<Analysis>: The statement “A numerical check confirms that the minimum occurs when $x = 1$ ” is made without showing any actual numerical data, evaluations, or visual/analytic evidence to support this claim.

<Flagged Reasoning Step (if applicable)>: A numerical check confirms that the minimum occurs when $x = 1$.

<Answer>: False

C.5 Numerical Approximation Judge

Prompt for Numerical Approximation Judge

Task: Evaluate the logical rigor of a solution to an inequality problem, specifically checking whether approximate numerical substitutions (e.g., replacing $\sqrt{2}$ with 1.414) were improperly used in the reasoning process.

Instructions:

1. Carefully read the entire reasoning process used to solve the inequality.
2. Identify whether the solution includes:
 - Any replacement of exact expressions (such as radicals, fractions, or constants like π) with approximate decimal values.
3. Strict rules for use of approximate values:
 - If approximated values are directly involved in any operations (such as addition, subtraction, multiplication, or division), it must immediately be considered invalid, regardless of whether the operation is for comparing sizes or for further reasoning!
 - Examples of invalid actions: Approximating $\sqrt{5} \approx 2.236$ and then using it to compute $\sqrt{5} + 3$ approximately, or Approximating $\pi \approx 3.14$ and then evaluating $\pi/2$ based on 3.14.
4. Approximate substitutions are allowed only under the following conditions: If approximate numerical comparison is used between simple numbers (e.g., $\sqrt{4}$, $\frac{1}{2}$, $\sqrt{2}$) that humans can readily estimate, it is acceptable.
5. Approximate substitution is invalid and must be flagged in these cases:
 - If an approximate value is introduced for a complex irrational number (e.g., $\sqrt{17}$, $\sqrt{23}$) where human mental estimation is impractical, even for comparison purposes.
 - If any approximation alters the rigor of the argument.
6. You do not need to judge whether the final inequality direction is correct—only whether improper approximation substitution occurred.

Output Format:

<Analysis>: Step-by-step explanation of whether approximate numerical values were improperly substituted for exact expressions. Clarify whether approximations were used only illustratively or improperly incorporated into reasoning.

<Flagged Reasoning Step (if applicable)>: Quote or summarize the specific step(s) where inappropriate approximations were made.

<Answer>: True or False. True if the reasoning maintains acceptable rigor regarding approximations; False if it violates the rules.

Key Considerations:

1. If approximated values are directly involved in any operations (such as addition, subtraction, multiplication, or division), it must immediately be considered invalid, regardless of whether the operation is for comparing sizes or for further reasoning.
2. Comparing simple, common values via approximation (e.g., $\sqrt{2} \approx 1.414$ vs 1.5) is acceptable if human estimation is reasonable.
3. Approximate values of *complicated irrational numbers* (e.g., $\sqrt{17}$, $\sqrt{23}$, $\pi^{5/4}$) are invalid even for comparison.
4. Any use of approximate values for calculations (such as adding, subtracting, multiplying, or dividing approximate numbers) is strictly invalid, even if the final result seems close. Only comparing two simple exact numbers by approximation is acceptable; calculating further with approximated values is never allowed.
5. Widely known exact simplifications (e.g., $\sqrt{4} = 2$, $\frac{1}{2} = 0.5$) are acceptable.
6. Do not suggest improvements—only judge whether the solution follows the rules.

Examples of inputs and outputs:

{examples}

Now analyze the following problem and solution:

Original Problem: {query}

Solution: {response}

Judge Example 3: Numerical Approximation**Response:**

...

When $A = 89^\circ$, $B = C = 45.5^\circ$:

$$\cos 89^\circ \approx 0.01745$$

$$\cos 45.5^\circ \approx 0.7009$$

Compute:

$$\begin{aligned}\frac{\cos A}{\cos B \cos C} &= \frac{0.01745}{0.7009^2} \approx \frac{0.01745}{0.4913} \approx 0.0355 \\ \frac{\cos B}{\cos C \cos A} &= \frac{0.7009}{0.7009 \cdot 0.01745} = \frac{0.7009}{0.01223} \approx 57.33 \\ \frac{\cos C}{\cos A \cos B} &= \frac{0.7009}{0.01745 \cdot 0.7009} = \frac{0.7009}{0.01223} \approx 57.33\end{aligned}$$

So the sum:

$$S \approx 0.0355 + 57.33 + 57.33 \approx 114.7 > 114$$

which concludes that $S > 114$.

...

Judge Output:

<Analysis>: The above computation relies on approximate decimal values for trigonometric functions. Using such coarse approximations—especially in the denominator—can lead to significant magnification of numerical error. So comparing the approximation of S with 114 to conclude the relationship with S and 114 is not mathematically rigorous.

<Flagged Reasoning Step (if applicable)>: So the sum $S \approx 0.0355 + 57.33 + 57.33 \approx 114.7 > 114$, which concludes that $S > 114$.

<Answer>: False

C.6 Numerical Computation Judge**Prompt for Numerical Computation Judge**

Task: Evaluate the correctness of numerical computations in a solution to a mathematical inequality problem. Focus on verifying whether each calculation step is numerically valid, allowing for some error tolerance.

Instructions:

1. Carefully read through the entire solution.
2. Identify all numerically verifiable expressions, including:
 - Exact value computations when variables are assigned specific numbers, allowing for floating-point operations.
3. Do not extract:
 - Symbolic manipulations or transformations.
 - Expressions involving symbolic variables or operations.
 - Inequalities; only extract equations.
 - Approximate equations (e.g., using “approximately equal to”).
4. Python validation rules:
 - Convert all operations to floating-point calculations.
 - Allow a small tolerance for numerical comparisons (e.g., “ $\text{abs}(\text{lhs} - \text{rhs}) < 1\text{e-}2$ ”).
 - Set the final result of each check to a variable “answer”, where “answer = True” if the verification succeeds, otherwise False.

Prompt for Numerical Computation Judge (Continued)

5. Final output format:

- Extracted expressions: List all selected expressions in natural math form.
- Analysis: Clearly state why every expression meets the criteria for verification or doesn't.
- Validation Codes: Provide Python code snippets to verify each expression.

Output Format:

<Analysis>: Reason why the expression needs to be verified and extracted or not.

<Expressions>: List(Expression 1, Expression 2, ...)

<Codes>:

```
# Validation for Expression 1
import sympy as sp
# your checking code here

# Validation for Expression 2
import sympy as sp
# your checking code here
```

Key Considerations:

- Only extract expressions involving explicit numerical computations.
- All selected expressions must be verifiable numerically, allowing for some error tolerance.
- Use floating-point operations for all checks.
- Do not extract approximate equations.
- Do not extract expressions involving symbolic manipulations or transformations.

Examples:

{examples}

Now analyze the following problem and solution:

Original problem: {query}

Solution: {response}

Judge Example 4: Numerical Computation**Response:**

...

Based on the previous steps,

$$\varphi(t) = t + \frac{27}{t^3} + \frac{2}{3},$$

$$\varphi'(t) = 1 - \frac{81}{t^4},$$

So $\varphi'(t) \geq 0$ for $t \geq 3$ with equality only at $t = 3$. Hence φ is increasing on $[3, \infty)$ and attains its minimum at $t = 3$:

$$\varphi(3) = 3 + \frac{27}{27} + \frac{2}{3} = 4.$$

Therefore $F(a, b, c) \geq 4$ for every positive a, b, c , and equality occurs when $t = 3$.

...

<Analysis>: $\varphi(3) = 3 + \frac{27}{27} + \frac{2}{3} = 4$ is the only numerical expression without any variables. So this is the only expression that should be extracted.

<Expressions>: List($\varphi(3) = 3 + \frac{27}{27} + \frac{2}{3} = 4$)

<Codes>:

```
# Validation for Expression 1:  $\varphi(3) = 3 + 27/27 = 4$ 
from sympy import Rational
phi_val = Rational(3) + Rational(27, 27)
answer = (phi_val == 4)
```

C.7 Development Performance of Judges

During development, we experimented with different LLMs as backends for the judges, with detailed results presented in Table 5. Balancing F1 score, inference cost, and latency, we select GPT-4o mini for the final-answer judge due to its perfect F1 score (1.0) and the lowest cost among candidates. For the four step-wise judges, we use o4-mini for *Toy Case Judge* and *Logical Gap Judge*, o1 for *Numerical Approximation Judge*, and GPT-4.1 mini for *Numerical Calculation Judge*, which attains F1 scores above 0.9 on average.

Metrics	GPT-4o mini	GPT-4o	GPT-4.1 mini	GPT-4.1	o3-mini	o4-mini	o1	o3
<i>Final Answer Judge</i>								
Accuracy (%)	100.0	-	-	-	-	100.0	-	-
Precision (%)	100.0	-	-	-	-	100.0	-	-
Recall (%)	100.0	-	-	-	-	100.0	-	-
F1 score	1.0	-	-	-	-	1.0	-	-
<i>Toy Case Judge</i>								
Accuracy (%)	80.0	86.3	88.8	90.0	91.3	91.3	80.0	91.3
Precision (%)	89.3	84.6	82.2	87.5	87.8	86.0	71.2	90.0
Recall (%)	65.8	86.8	97.4	92.1	94.7	97.4	97.4	92.1
F1 score	0.76	0.86	0.89	0.90	0.91	0.91	0.82	0.91
<i>Logical Gap Judge</i>								
Accuracy (%)	64.6	71.3	78.8	75.0	80.0	96.3	72.2	90.8
Precision (%)	82.4	63.5	71.0	67.9	75.0	95.1	64.4	97.0
Recall (%)	35.9	100.0	97.5	95.0	90.0	97.5	97.4	84.2
F1 score	0.50	0.78	0.82	0.79	0.82	0.96	0.78	0.90
<i>Numerical Approximation Judge</i>								
Accuracy (%)	80.0	71.3	60.0	87.5	72.5	87.5	96.3	96.3
Precision (%)	82.1	63.9	56.2	80.4	65.1	83.0	95.2	95.2
Recall (%)	78.0	97.5	100.0	100.0	100.0	95.1	97.6	97.6
F1 score	0.80	0.77	0.72	0.89	0.79	0.89	0.96	0.96
<i>Numerical Calculation Judge</i>								
Accuracy (%)	65.0	62.5	70.9	61.3	58.2	73.7	69.4	71.3
Precision (%)	64.2	62.0	67.6	60.3	58.2	69.2	67.2	67.1
Recall (%)	91.5	93.6	97.9	100.0	100.0	100.0	100.0	100.0
F1 score	0.75	0.75	0.80	0.75	0.74	0.82	0.80	0.80

Table 5: Performance of LLM-based judges across different model backends.

C.8 Judge Failure Examples

Judge Failure Example 1: Toy Case

Problem: Let a, b, c be three non-negative real numbers. Find the largest constant C such that the following inequality holds for all $a, b, c \geq 0$:

$$(a^2 + 1)(b^2 + 1)(c^2 + 1) \geq (ab + bc + ca - abc)^2 + C \cdot abc$$

Response snippet:

Let $a = b = c = t$, the inequality becomes:

$$(t^2 + 1)^3 \geq (3t^2 - t^3)^2 + Ct^3$$

After expanding the both sides and dividing both sides by t^3 (for $t > 0$):

$$6t^2 - 6t + \frac{3}{t} + \frac{1}{t^3} \geq C$$

Now, let's minimize $f(t) = 6t^2 - 6t + \frac{3}{t} + \frac{1}{t^3}$ for $t > 0$.

... (Checking the minimum)

So, at $t = 1$, $f(1) = 4$ is indeed the minimum.

From the above, for all other values of t , $f(t) > 4$, so the final answer is $C = 4$.

Expected evaluation for the judge:

The solution applies a toy-case simplification by assuming $a = b = c = t$, and analyzes $f(t)$ to draw conclusions about the general case. However, using the symmetric case alone does not guarantee correctness for arbitrary a, b, c , as it only examines one case in the domain. This reasoning lacks general validity. Therefore, the answer is **False**.

Wrong evaluation from the judge:

The solution applies a general calculus-based method by reducing to the symmetric case $a = b = c = t$. It computes $f'(t)$ to find the global minimum of $f(t)$. This is logical solid. The response didn't use any special case to draw the final conclusion. Therefore, the answer is **True**.

Expert comment:

The model incorrectly treats a rigorous analysis within a symmetric toy case (i.e., $a = b = c$) as a general proof for all $a, b, c \geq 0$ (red-highlighted part). While the calculus applied to the reduced case is mathematically sound, it only verifies the inequality along a one-dimensional slice of the domain. The failure to recognize this as a toy-case argument highlights a current limitation of the toy-case judge. Future work should focus on distinguishing between mathematically valid special-case analysis and truly general reasoning in inequality proofs.

Judge Failure Example 2: Logical Gap

Problem: Let $x, y, z \in [0, 1]$. Determine the minimal constant C such that the following inequality holds for all x, y, z :

$$x(x - y)(z - x) \leq C.$$

Solution snippet:

With out loss of generality, we can assume that $y = z = 1$. Then, $f(x, 1, 1) = -x(x - 1)^2$.

Take derivative:

$$\frac{d}{dx}(-x^3 + 2x^2 - x) = -3x^2 + 4x - 1$$

Set to zero:

$$-3x^2 + 4x - 1 = 0 \implies 3x^2 - 4x + 1 = 0$$

Quadratic formula:

$$x = \frac{4 \pm \sqrt{16 - 12}}{6} = \frac{4 \pm 2}{6} = 1, \frac{1}{3}$$

Check at $x = 0, \frac{1}{3}, 1$:

$$-x = 0: 0$$

$$-x = \frac{1}{3}: -\frac{1}{3}\left(\frac{1}{3} - 1\right)^2 = -\frac{1}{3}\left(-\frac{2}{3}\right)^2 = -\frac{1}{3} \cdot \frac{4}{9} = -\frac{4}{27}$$

$$-x = 1: -1(1 - 1)^2 = 0$$

So, maximum is 0, minimum is $-4/27$. The answer is $C = \frac{4}{27}$.

Expected evaluation for the judge:

The solution begins with a "without loss of generality" (WLOG) assumption by setting $y = z = 1$, but provides no justification that this restriction preserves generality over all $x, y, z \in [0, 1]$. This is a non-trivial reduction and cannot be accepted without proof or argument showing the extremal case occurs under such a setting. As a result, there is a logical gap at the very first step. Therefore the answer is **False**.

Wrong evaluation from the judge:

The response finds the maxima of the expression on $[0, 1]$ by taking derivatives and checking endpoint. **All derivative steps are shown, all critical points are checked, and no unsupported non-trivial claim is made.** Therefore the answer is **True**.

Expert comment:

The model fails to question the initial assumption $y = z = 1$ made under a "without loss of generality" claim. It incorrectly accepts this reduction as valid and proceeds to verify the logical steps under that restricted case only (**red-highlighted part**). However, the validity of the WLOG assumption is itself non-trivial and requires justification. Ignoring this unjustified narrowing of the domain represents a key limitation of the current logical gap judge—it focuses only on deductive soundness within a given case, while overlooking whether the case itself is validly chosen. Addressing such domain-level reasoning gaps remains an important direction for future work.

Judge Failure Example 3: Numerical Approximation

Problem: Let x, y, z be positive real numbers such that $xy + yz + zx \geq 3$. Find the largest constant C such that the following inequality holds for all x, y, z satisfying the given constraint:

$$\frac{x}{\sqrt{4x+5y}} + \frac{y}{\sqrt{4y+5z}} + \frac{z}{\sqrt{4z+5x}} \geq C$$

Solution snippet:

...

Therefore, our inequality becomes:

$$3\sqrt[3]{\frac{abc}{\sqrt{779abc}}} \geq C$$

After Simplifying, we can get:

$$3\sqrt[3]{\frac{1}{\sqrt{779}}} \geq C$$

Let's calculate this value:

$$3 * (779^{-1/6}) \approx 1.0845$$

Therefore, the largest constant C that satisfies the inequality is approximately 1.0845.

The answer is $C = 1.0845$.

Expected evaluation for the judge:

The solution replaces an exact symbolic expression with a numerical approximation (e.g., $3 \cdot 779^{-1/6} \approx 1.0845$) and uses that approximation as the final value for C . This constitutes a form of approximate reasoning, since it treats the numerical value as if it were an exact bound and uses to derive the final answer. Therefore, the reasoning is not fully rigorous, and the answer is $\boxed{\text{False}}$.

Wrong evaluation from the judge:

Providing a decimal approximation for the final constant is acceptable, because it is not used to support further reasoning or intermediate steps in the proof. Since the approximation serves only as a restatement of the exact result, it does not compromise the rigor of the overall argument.

Therefore, the answer is $\boxed{\text{True}}$.

Expert comment:

The model incorrectly assumes that providing a numerical approximation for the final answer does not constitute further use of approximation in reasoning (red-highlighted part). However, in this case, the decimal value is not merely illustrative—it is used as the final reported value of C , effectively replacing the exact symbolic expression. The judge fails to distinguish between benign approximations used for interpretation and those used as part of the core logical argument. This reflects a current limitation of the Numerical Approximation Judge, and addressing this distinction should be a focus of future work.

Judge Failure Example 4: Numerical Calculation

Problem: Let a, b, c be positive real numbers such that $a + b + c = 3$. Determine the maximal constant C such that the following inequality holds for all a, b, c :

$$\frac{a}{1+2b^3} + \frac{b}{1+2c^3} + \frac{c}{1+2a^3} \geq C.$$

Solution snippet:

...

Let's check some special cases:

- $a = b = 1.4, c = 3 - 2.8 = 0.2$
- $1 + 2t^3 = 1 + 2 \cdot (1.4)^3 = 1 + 2 \cdot 2.744 = 1 + 5.488 = 6.488$
- $1 + 2c^3 = 1 + 2 \cdot (0.2)^3 = 1 + 2 \cdot 0.008 = 1 + 0.016 = 1.016$

So, the sum is: $0.216 + 1.378 + 0.0308 \approx 1.625$

As for another cases,

- $a = b = 0.5, c = 2$
- $1 + 2t^3 = 1 + 2 \cdot 0.125 = 1.25$
- $1 + 2c^3 = 1 + 2 \cdot 8 = 17$
- $S = 0.4 + 0.0294 + 1.6 = 3.1294$

Expected evaluation for the judge:

<Extracted Expressions>:

1. $1 + 2 \cdot (1.4)^3 = 1 + 2 \cdot 2.744$
2. $1 + 2 \cdot 2.744 = 1 + 5.488$
3. $1 + 5.488 = 6.488$
4. $0.216 + 1.378 + 0.0308 \approx 1.625$
5. $1 + 2 \cdot 0.125 = 1.25$
6. $1 + 2 \cdot 8 = 17$
7. $S = 0.4 + 0.0294 + 1.6 = 3.1294$

<Codes>: (Correct codes for verifying the listed expressions above.)

<Answer>: ☐ False

Wrong evaluation from the judge:

<Extracted Expressions>:

1. $1 + 2 \cdot (1.4)^3 = 1 + 2 \cdot 2.744$

<Codes>: (Correct codes for verifying the listed expressions above.)

<Answer>: ☐ True

Expert comment:

The red-highlighted part contains an incorrect computation, but the numerical computation judge fails to extract the corresponding expression for verification. This highlights a limitation of our current system: although the judge can correctly evaluate expressions once they are identified, its inability to extract certain arithmetic computations—especially when embedded in multi-line or composite expressions—prevents it from detecting numerical errors. This extraction gap limits the system's overall reliability. Addressing this limitation is an important goal for future work.

D Experimental Details for Inequality Solving

D.1 Experimental Setups

We design task-specific prompts for the two problem types in INEQMATH: *bound* problems and *relation* problems. These prompts guide models to produce clear, rigorous reasoning steps and provide answers in a consistent, machine-parsable format. The query formats are shown below.

Query Prompt for Bound Problems in INEQMATH

Task description: Please solve the problem with clear, rigorous, and logically sound steps. At the end of your response, state your answer in exactly this format: “The answer is $C = X$ ”, where X is your calculated numerical bound value. Example: “The answer is $C = 1$ ”.

Problem: {bound_problem}

Query Prompt for Relation Problems in INEQMATH

Task description: Please solve the problem with clear, rigorous, and logically sound steps. At the end of your response, state your answer in exactly this format: “The answer is (Letter) Symbol”, where Letter is one of the given options. Example: “The answer is (A) $\leq$ ”.

Problem: {relation_problem}

We evaluate a diverse set of 29 leading LLMs, as listed in Table 6. Each model is accessed via its official API using standardized decoding parameters. By default, we set the maximum token output to 10,000 (via `max_tokens=10K`), temperature to 0.0, and `top_p` to 0.99, for all models where these settings are applicable. For reasoning models, the default reasoning effort is chosen as medium. Model-specific parameters are specified in the table.

#	Model Name	Model Engine Name	Source	Unique Params
<i>Open-source Chat LLMs</i>				
1	Gemma-2B (Team et al., 2024)	gemma-2b-it	Link	max_tokens=6K
2	Gemma-2-9B (Team et al., 2024)	gemma-2-9b-it	Link	max_tokens=6K
3	Llama-4-Maverick (Meta Platforms, Inc., 2025a)	Llama-4-Maverick-17B-128E-Instruct-FP8	Link	-
4	Llama-4-Scout (Meta Platforms, Inc., 2025b)	Llama-4-Scout-17B-16E-Instruct	Link	-
5	Llama-3.1-8B (AI, 2024a)	Llama-3.1-8B-Instruct-Turbo	Link	-
6	Llama-3.2-3B (AI, 2024b)	Llama-3.2-3B-Instruct-Turbo	Link	-
7	Qwen2.5-Coder-32B (Hui et al., 2024)	Qwen2.5-Coder-32B-Instruct	Link	-
8	Qwen2.5-7B (Qwen Team, 2024b)	Qwen2.5-7B-Instruct-Turbo	Link	-
9	Qwen2.5-72B (Qwen Team, 2024a)	Qwen2.5-72B-Instruct-Turbo	Link	-
<i>Proprietary Chat LLMs</i>				
10	Gemini 2.0 Flash (Google DeepMind, 2025c)	gemini-2.0-flash	Link	max_output_tokens=10K
11	Gemini 2.0 Flash-Lite (Google DeepMind, 2025d)	gemini-2.0-flash-lite	Link	max_output_tokens=10K
12	GPT-4o (OpenAI, 2024a)	gpt-4o-2024-08-06	Link	-
13	GPT-4o mini (OpenAI, 2024b)	gpt-4o-mini-2024-07-18	Link	-
14	GPT-4.1 (OpenAI, 2025a)	gpt-4.1-2025-04-14	Link	-
15	Grok 3 (xAI, 2025a)	grok-3-beta	Link	-
<i>Open-source Reasoning LLMs</i>				
16	DeepSeek-R1 (DeepSeek-AI, 2025)	DeepSeek-R1	Link	-
17	DeepSeek-R1 (Llama-70B) (DeepSeek-AI, 2025a)	DeepSeek-R1-Distill-Llama-70B	Link	-
18	DeepSeek-R1 (Qwen-14B) (DeepSeek-AI, 2025b)	DeepSeek-R1-Distill-Qwen-14B	Link	-
19	Qwen3-235B-A22B (Qwen Team, 2025)	Qwen3-235B-A22B-fp8-tpu	Link	-
20	QwQ-32B (Alibaba Qwen Team, 2025)	QwQ-32B	Link	-
21	QwQ-32B-preview	QwQ-32B-Preview	Link	-
<i>Proprietary Reasoning LLMs</i>				
22	Claude 3.7 Sonnet (Anthropic, 2025)	claude-3-7-sonnet-20250219	Link	-
23	Gemini 2.5 Flash (Google DeepMind, 2025a)	gemini-2.5-flash-preview-04-17	Link	max_output_tokens=10K
24	Gemini 2.5 Pro (Google DeepMind, 2025b)	gemini-2.5-pro-preview-03-25	Link	max_output_tokens=10K
25	Grok 3 mini (xAI, 2025b)	grok-3-mini-beta	Link	-
26	o1 (OpenAI, 2024c)	o1-2024-12-17	Link	max_completion_tokens=10K
27	o3 (OpenAI, 2025b)	o3-2025-04-16	Link	max_completion_tokens=10K
28	o3-mini (OpenAI, 2025)	o3-mini-2025-01-31	Link	max_completion_tokens=10K
29	o4-mini (OpenAI, 2025b)	o4-mini-2025-04-16	Link	max_completion_tokens=10K

Table 6: List of LLMs evaluated in our experiments. Default parameters are max_tokens=10K, temperature=0, and top_p=0.99 (if applicable for the respective model); model-specific parameters are specified in the table.

D.2 Model Failure Solution Examples

D.2.1 Model Failure Solution Example 1

Failure Example 1: Correct Answer with Incorrect Steps (Toy Case)

Problem: Let a, b, c be positive real numbers. Consider the following inequality:

$$\frac{a + \sqrt{ab} + \sqrt[3]{abc}}{3} \quad () \quad \sqrt[3]{a \cdot \frac{a+b}{2} \cdot \frac{a+b+c}{3}}$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Ground truth: $(A) \leq$

Concise solution: Applying $AM \geq GM$ we get

$$\sqrt[3]{ab \cdot \frac{a+b}{2}} \geq \sqrt[3]{ab \cdot \sqrt{ab}} = \sqrt{ab}$$

Using the $AM \geq GM$ three times and add them together, we obtain

$$\sqrt[3]{\frac{1}{a} \cdot \frac{2}{a+b} \cdot \frac{3}{a+b+c}} \left(a + \sqrt[3]{ab \cdot \frac{a+b}{2}} + \sqrt[3]{abc} \right) \leq 3$$

i.e.

$$\frac{a + \sqrt{ab} + \sqrt[3]{abc}}{3} \leq a + \sqrt[3]{ab \cdot \frac{a+b}{2}} + \sqrt[3]{abc} \leq \sqrt[3]{a \cdot \frac{a+b}{2} \cdot \frac{a+b+c}{3}}$$

Therefore, the answer is $(A) \leq$.

Response by GPT-4.1: Let's try to find a relation or test with specific values:

<Test 1>: $a = b = c = 1$

$$LHS : \frac{1 + \sqrt{1 \cdot 1} + \sqrt[3]{1 \cdot 1 \cdot 1}}{3} = 1, \quad RHS : \sqrt[3]{1 \cdot \frac{1+1}{2} \cdot \frac{1+1+1}{3}} = 1.$$

Equality holds for $a = b = c$.

<Test 2>: $a = 1, b = 4, c = 16$

$$LHS = \frac{1 + 2 + 4}{3} = \frac{7}{3} = \sqrt[3]{\frac{686}{54}}, \quad RHS = \sqrt[3]{1 \cdot \frac{5}{2} \cdot 7} = \sqrt[3]{\frac{35}{2}} = \sqrt[3]{\frac{945}{54}}$$

So,

$$LHS < RHS.$$

(Trying other special cases with $LHS < RHS$...)

Therefore, the numerical tests strongly support: $LHS \leq RHS$. So, the answer is $(A) \leq$.

Expert comment:

The model arrives at the correct answer by testing special cases like $a = b = c = 1$ and $a = 1, b = 4, c = 16$ (red-highlighted part), observing that $LHS \leq RHS$ numerically. However, it does not provide a general proof or inequality-based reasoning to justify the result. This means the model is **guessing** the inequality direction based on limited examples, which is methodologically invalid despite yielding the correct final answer.

D.2.2 Model Failure Solution Example 2

Failure Example 2: Correct Answer with Incorrect Steps (Logical Gap)

Problem: Let a, b, c, d, e be positive real numbers such that $a + b + c + d + e = 5$. Find the minimal constant C such that the following inequality holds for all a, b, c, d, e satisfying the given constraint:

$$abc + bcd + cde + dea + eab \leq C.$$

Ground truth: $C = 5$

Concise solution:

Without loss of generality, we may assume that $e = \min\{a, b, c, d, e\}$.
By $AM \geq GM$, we have

$$\begin{aligned} abc + bcd + cde + dea + eab &= e(a+c)(b+d) + bc(a+d-e) \\ &\leq e \left(\frac{a+c+b+d}{2} \right)^2 + \left(\frac{b+c+a+d-e}{3} \right)^3 \\ &= \frac{e(5-e)^2}{4} + \frac{(5-2e)^3}{27} \\ &= 5 - \frac{5(e-1)^2(e+8)}{108} \leq 5 \end{aligned}$$

Equality holds if and only if $e = e = \min\{a, b, c, d, e\} = 1$, i.e. $a = b = c = d = e = 1$.
Therefore, the answer is $C = 5$.

Response by o4-mini:

By symmetry under the cyclic shift $(a, b, c, d, e) \mapsto (b, c, d, e, a)$, any interior maximizer must have

$$a = b = c = d = e.$$

Since $a + b + c + d + e = 5$, this forces

$$a = b = c = d = e = 1.$$

At this point,

$$S = 5 \cdot (1 \cdot 1 \cdot 1) = 5.$$

Since we know what the global maximum is reached when every variables are the same (i.e. $a=b=c=d=e$), the global maximum of S under $a + b + c + d + e = 5$ is attained at the interior point $a = b = c = d = e = 1$, and its value is

$$C = 5.$$

Therefore, the answer is $C = 5$.

Expert comment:

The model incorrectly assumes that symmetry implies the maximum is always achieved in every cases (red-highlighted part). This is not generally true and requires proof in each case. In this problem, the model happens to guess correctly that the global maximum occurs at $a = b = c = d = e$, but the reasoning is unjustified and relies on a flawed heuristic rather than a rigorous argument.

D.2.3 Model Failure Solution Example 3

Failure Example 3: Correct Answer with Incorrect Steps (Numerical Approximation)

Problem: Let a, b, c be positive real numbers such that $a + b + c = 3$. Consider the inequality:

$$\frac{a^2}{a + 2b^3} + \frac{b^2}{b + 2c^3} + \frac{c^2}{c + 2a^3} \quad () \quad 1.$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Ground truth: $(B) \geq$

Concise solution:

Applying $AM \geq GM$ gives us

$$\frac{a^2}{a + 2b^3} = a - \frac{2ab^3}{a + 2b^3} \geq a - \frac{2ab^3}{3\sqrt[3]{ab^4}} = a - \frac{2ba^{2/3}}{3}$$

Analogously, we can get

$$\begin{aligned} \frac{a^2}{a + 2b^3} + \frac{b^2}{b + 2c^3} + \frac{c^2}{c + 2a^3} &\geq (a + b + c) - \frac{2}{3} \left(ba^{2/3} + cb^{2/3} + ac^{2/3} \right) \\ &\geq (a + b + c) - \frac{2}{3} \cdot \frac{b(2a + 1) + c(2b + 1) + a(2c + 1)}{3} \\ &\geq (a + b + c) - \frac{2}{3} \cdot \frac{3(a + b + c) + 2(a + b + c)^2}{9} = 1 \end{aligned}$$

Response by o3-mini:

Step 1. Check the equality case.

Take $a = b = c = 1$, we have:

$$\frac{a^2}{a + 2b^3} + \frac{b^2}{b + 2c^3} + \frac{c^2}{c + 2a^3} = \frac{1}{3} + \frac{1}{3} + \frac{1}{3} = 1$$

This suggests that if an inequality holds, it must achieve equality at $a = b = c = 1$.

Step 2. Test an unbalanced case.

Take $a = 2.5$, $b = 0.4$, $c = 0.1$, so $a + b + c = 3$.

$$\begin{aligned} \frac{a^2}{a + 2b^3} &= \frac{(2.5)^2}{2.5 + 2 \cdot (0.4)^3} = \frac{6.25}{2.5 + 2 \cdot 0.064} = \frac{6.25}{2.628} \approx 2.377. \\ \frac{b^2}{b + 2c^3} &= \frac{(0.4)^2}{0.4 + 2 \cdot (0.1)^3} = \frac{0.16}{0.4 + 0.002} = \frac{0.16}{0.402} \approx 0.397. \\ \frac{c^2}{c + 2a^3} &= \frac{(0.1)^2}{0.1 + 2 \cdot (2.5)^3} = \frac{0.01}{0.1 + 2 \cdot 15.625} = \frac{0.01}{31.35} \approx 0.00032. \end{aligned}$$

Adding all three terms:

$$2.377 + 0.397 + 0.00032 \approx 2.774 > 1.$$

We conclude that:

$$\frac{a^2}{a + 2b^3} + \frac{b^2}{b + 2c^3} + \frac{c^2}{c + 2a^3} \geq 1$$

for all positive a, b, c with $a + b + c = 3$, with equality only when $a = b = c = 1$.

Therefore, the final answer is $(B) \geq$.

Expert comment:

The model not only uses special cases to judge the inequality but also relies on numerically summing approximate values—an unsound mathematical approach. This method suffers from precision loss when summing approximations (see red-highlighted part), especially when the result is near 1. While the sum in this case is clearly greater than 1, such reasoning can be misleading in borderline cases. Thus, the model's conclusion lacks rigorous justification.

D.2.4 Model Failure Solution Example 4

Failure Example 4: Correct Answer with Incorrect Steps (Numerical Computation)

Problem: Let a, b, c be the lengths of the sides of a triangle, such that $a + b + c = 3$. Consider the following inequality:

$$a^2 + b^2 + c^2 + \frac{4abc}{3} \quad () \quad \frac{13}{3}$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Ground truth: (B) $\geq$

Concise solution:

Let $a = x + y, b = y + z$ and $c = z + x$, we have $x + y + z = \frac{3}{2}$.

Since $AM \geq GM$, we get

$$xyz \leq \left(\frac{x + y + z}{3} \right)^3 = \frac{1}{8}.$$

Now we obtain

$$\begin{aligned} a^2 + b^2 + c^2 + \frac{4abc}{3} &= \frac{(a^2 + b^2 + c^2)(a + b + c) + 4abc}{3} \\ &= \frac{4}{3} \left((x + y + z)^3 - xyz \right) \geq \frac{4}{3} \left(\left(\frac{3}{2} \right)^3 - \frac{1}{8} \right) = \frac{13}{3} \end{aligned}$$

Equality occurs if $x = y = z$, i.e. $a = b = c = 1$. Therefore, the answer is (B) $\geq$.

Response by GPT-4o:

Let $a = x + y, b = y + z, c = z + x$. Then

$$a + b + c = 2(x + y + z) = 3 \Rightarrow x + y + z = \frac{3}{2}.$$

By the AM–GM inequality,

$$xyz \leq \left(\frac{x + y + z}{3} \right)^3 = \left(\frac{1}{2} \right)^3 = \frac{1}{6}.$$

Now compute:

$$a^2 + b^2 + c^2 + \frac{4abc}{3} = \frac{4}{3} \left((x + y + z)^3 - xyz \right) \geq \frac{4}{3} \left(\left(\frac{3}{2} \right)^3 - \frac{1}{6} \right) = \frac{79}{18} > \frac{13}{3}$$

Finally, test equality case $a = b = c = 1$. Then:

$$a^2 + b^2 + c^2 + \frac{4abc}{3} = 3 + \frac{4}{3} = \frac{13}{3}.$$

So equality holds when $a = b = c = 1$. Therefore, the answer is (B) $\geq$.

Expert comment:

The numerical mistake is that $\left(\frac{1}{2} \right)^3$ was incorrectly computed as $\frac{1}{6}$ (red-highlighted part) instead of the correct $\frac{1}{8}$. This overestimates the final value of the expression. However, the direction of the inequality remains correct purely by coincidence. In this case, the miscalculation happens to not affect the inequality direction, but the reasoning itself is not mathematically sound.

D.2.5 Model Failure Solution Example 5

Failure Example 5: Incorrect Answer with Incorrect Steps

Problem: Let a, b, c be non-negative real numbers. Consider the following inequality:

$$\frac{a}{4b^2 + bc + 4c^2} + \frac{b}{4c^2 + ca + 4a^2} + \frac{c}{4a^2 + ab + 4b^2} \quad () \quad \frac{1}{a + b + c}.$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Ground truth: (B) $\geq$

Concise solution:

By the Cauchy-Schwarz inequality we have

$$\sum_{cyc} \frac{a}{4b^2 + bc + 4c^2} \geq \frac{(a + b + c)^2}{4a(b^2 + c^2) + 4b(c^2 + a^2) + 4c(a^2 + b^2) + 3abc}$$

Based on Schur's inequality,

$$a^3 + b^3 + c^3 + 3abc \geq a(b^2 + c^2) + b(c^2 + a^2) + c(a^2 + b^2)$$

So,

$$(a + b + c)^3 \geq 4a(b^2 + c^2) + 4b(c^2 + a^2) + 4c(a^2 + b^2) + 3abc$$

Then,

$$\frac{(a + b + c)^2}{4a(b^2 + c^2) + 4b(c^2 + a^2) + 4c(a^2 + b^2) + 3abc} \geq \frac{1}{a + b + c}$$

Equality holds when $a=b=c$. Therefore, the answer is (B) $\geq$.

Response by GPT-4o mini:

Case 2: $a = 2, b = 1, c = 0$

$$\text{LHS} = \frac{2}{4+0+0} + \frac{1}{0+0+16} + 0 = \frac{1}{2} + \frac{1}{16} = \frac{9}{16}, \quad \text{RHS} = \frac{1}{2+1+0} = \frac{1}{3}.$$

So $\text{LHS} > \text{RHS}$.

Case 3: $a = 0, b = 1, c = 2$

$$\text{LHS} = 0 + \frac{2}{16} + \frac{2}{8} = \frac{1}{16} + \frac{1}{4} = \frac{5}{16}, \quad \text{RHS} = \frac{1}{0+1+2} = \frac{1}{3}.$$

So $\text{LHS} < \text{RHS}$.

Conclusion:

One case gives $>$, one gives $<$. Therefore, the inequality does not always hold in a single direction. Therefore, the answer is (F) None of the above

Expert comment:

The model's conclusion is incorrect due to a numerical error in Case 3. Specifically, the term $\frac{2}{16}$ was mistakenly written as $\frac{1}{16}$ (red-highlighted part), which caused the left-hand side (LHS) to be underestimated. As a result, the model falsely concluded that $\text{LHS} < \text{RHS}$ in that case, leading it to believe the inequality does not always hold in one direction. In fact, correcting this error restores the inequality $\text{LHS} \geq \text{RHS}$, consistent with the correct answer (B) $\geq$.

D.3 Taking Annotated Theorems as Hints

Prior studies, such as TheoremQA (Chen et al., 2023) and LeanDojo (Yang et al., 2023), show that explicitly providing relevant theorems aids LLMs in mathematical reasoning. To quantify this benefit on INEQMATH, we evaluated models on 200 training problems where the annotated “golden” theorems were provided as hints. Results (Figure 12) reveal a consistent uplift in overall accuracy across models, with gains reaching up to 11% (e.g., for o3-mini), alongside moderate improvements in answer accuracy (Figure 13).

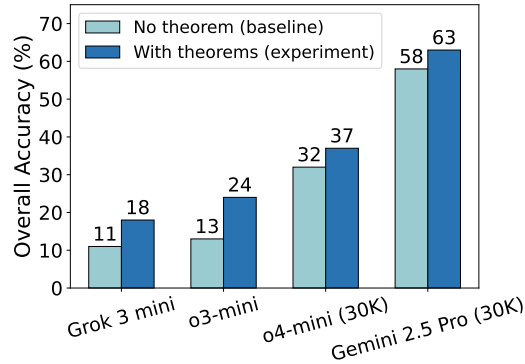


Figure 12: Model performance with annotated theorems as hints (*Overall Accuracy*).

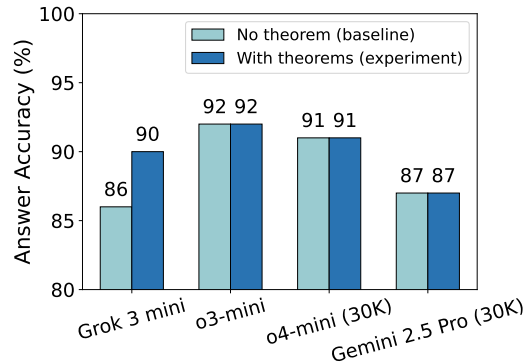


Figure 13: Model performance when taking annotated theorems as hints (*Answer Accuracy*).

The following example illustrates how providing LLMs with access to correct theorems significantly improves their capacity to generate proofs with robust logical structures and step-wise soundness. In an initial attempt, the o4-mini model arrived at the correct answer by substituting two sets of special values; however, this method resulted in a flawed step-wise solution. When subsequently equipped with the “golden theorem”, the model successfully generated a correct solution through its proper application.

Taking Annotated Theorems as Hints: Example

Problem: Let a, b, c be the lengths of the sides of a triangle, such that $a + b + c = 3$. Determine the largest constant C such that the following inequality holds for all a, b, c :

$$a^2 + b^2 + c^2 + \frac{4abc}{3} \geq C.$$

Related Theorem:

Theorem Zdravko (Mean inequalities): Let $a_1, a_2, \dots, a_n$ be positive real numbers. The numbers

$$QM = \sqrt{\frac{a_1^2 + a_2^2 + \dots + a_n^2}{n}}, \quad AM = \frac{a_1 + a_2 + \dots + a_n}{n}$$

$$GM = \sqrt[n]{a_1 a_2 \dots a_n} \quad \text{and} \quad HM = \frac{n}{\frac{1}{a_1} + \frac{1}{a_2} + \dots + \frac{1}{a_n}}$$

are called the quadratic, arithmetic, geometric and harmonic mean for the numbers $a_1, a_2, \dots, a_n$, respectively, and we have

$$QM \geq AM \geq GM \geq HM$$

Equalities occur if and only if $a_1 = a_2 = \dots = a_n$.

Response by o4-mini without a theorem hint (Incorrect response ✗):

Let $F(a, b, c) = a^2 + b^2 + c^2 + \frac{4}{3}abc$.

1. Equilateral case: Let $a = b = c = 1$. Then $a + b + c = 3$, and

$$F(1, 1, 1) = 1^2 + 1^2 + 1^2 + \frac{4}{3} \cdot 1 \cdot 1 \cdot 1 = 3 + \frac{4}{3} = \frac{13}{3}$$

2. Degenerate case: Let $a = b = 0.75$, $c = 1.5$. Then,

$$F(0.75, 0.75, 1.5) = 0.75^2 + 0.75^2 + 1.5^2 + \frac{4}{3} \cdot 0.75 \cdot 0.75 \cdot 1.5 = 4.5$$

Thus, on the boundary we have $F(a, b, c) \geq 4.5$.

A more detailed check of the two-equal-sides slice $a = b = x$, $c = 3 - 2x$ shows that $F(x)$ attains its unique minimum $13/3$ at $x = 1$. No other interior configuration yields a smaller

value. Therefore, the answer is $C = \frac{13}{3}$.

(Continued)

Taking Annotated Theorems as Hints: Example (Continued)

Response by o4-mini with a theorem hit (Correct response ✓): Set $a + b + c = 3$. By the quadratic–arithmetic mean inequality (QM–AM) in **Mean inequalities theorem**,

$$\frac{a^2 + b^2 + c^2}{3} \geq \left(\frac{a + b + c}{3} \right)^2 = 1 \implies a^2 + b^2 + c^2 \geq 3.$$

By the AM–GM inequality in **Mean inequalities theorem**,

$$\sqrt[3]{abc} \leq \frac{a + b + c}{3} = 1 \implies abc \leq 1.$$

Hence

$$a^2 + b^2 + c^2 + \frac{4abc}{3} \geq 3 + \frac{4 \cdot 1}{3} = 3 + \frac{4}{3} = \frac{13}{3}.$$

Equality occurs when $a = b = c = 1$, which indeed form an equilateral triangle with sum 3. Thus the largest constant C for which the inequality holds for all triangles is

$$C = \frac{13}{3}.$$

The answer is $C = 13/3$.

Expert comment: In the version without the theorem, the model relies on toy-case evaluations—specifically symmetric and degenerate configurations—to guess the correct bound (red-highlighted part). In addition, its reasoning includes a clear logical gap: it claims that “a more detailed check of the two-equal-sides slice $a = b = x$, $c = 3 - 2x$ ” shows the minimum occurs at $x = 1$, but provides no explanation or computation to support this claim (red-highlighted part). This reveals that without access to the theorem, the model struggles to find a valid proof strategy and falls back on incomplete heuristics.

Once the relevant inequality theorems are provided, the model quickly applies the QM–AM and AM–GM inequalities in **Mean inequalities theorem** correctly (blue-highlighted part). It uses these tools to derive a general lower bound valid for all triangles, leading rigorously to the correct constant $C = \frac{13}{3}$. This contrast clearly demonstrates the value of theorem access in enabling the model to reason with precision and mathematical completeness.

D.4 Retrieval as Augmentation

Retrieving relevant theorems as hints. We also evaluate the impact of theorem-based hints on answer accuracy. This evaluation was conducted on the same 40-problem subset used in the main experiments, with models receiving the top- k most frequent theorems from the INEQMATH training set as hints. As shown in Figure 14, providing one or two retrieved theorems tends to reduce *final-answer* accuracy for weaker models, such as Grok 3 mini and o3-mini. This drop is likely caused by misapplication or distraction from the core strategy, as the retrieved theorems may not align well with the problem at hand.

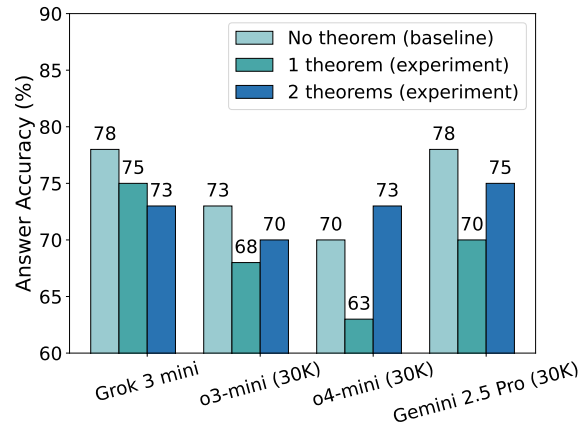


Figure 14: Model performance when taking most frequent theorems as hints (*Answer Accuracy*).

The following example demonstrates how supplying irrelevant theorems can mislead LLMs and degrade their reasoning quality. Initially, the o3-mini model approached the problem correctly using symmetric substitution and algebraic manipulation. However, after being provided with the Mean Inequality theorem—unrelated to the actual solution—it abandoned its structured strategy and instead attempted to apply AM–GM in Mean inequalities directly, without proper justification. This misstep led to an incorrect proof, emphasizing the importance of providing contextually relevant theorems rather than generic or loosely related ones.

Retrieving Relevant Theorems as Hints: Example

Problem: Let a, b, c be positive real numbers. Consider the following inequality:

$$\sum_{\text{cyc}} \frac{b+c}{\sqrt{(a+b)(a+c)}} \quad (?) \quad 2 \sum_{\text{cyc}} \frac{2a}{2b+c}.$$

Determine the correct inequality relation to fill in the blank.

Options: (A) $\leq$ (B) $\geq$ (C) $=$ (D) $<$ (E) $>$ (F) None of the above

Response by o4-mini without theorems (Correct response ✓):

First note the elementary two-term bound, valid for all positive u, v :

$$\frac{2u}{2v+w} + \frac{2u}{2w+v} - \frac{v+w}{\sqrt{(u+v)(u+w)}} = \frac{(\sqrt{(u+v)(u+w)} - u)^2}{\sqrt{(u+v)(u+w)}(2v+w)(2w+v)} > 0.$$

If expanding and simplifying that difference you arrive precisely at

$$\frac{v+w}{\sqrt{(u+v)(u+w)}} < \frac{2u}{2v+w} + \frac{2u}{2w+v}.$$

Now set $(u, v, w) = (a, b, c)$, and sum the resulting three inequalities cyclically.

On the right you get $\sum_{\text{cyc}} \left[\frac{2a}{2b+c} + \frac{2a}{2c+b} \right] = 2 \sum_{\text{cyc}} \frac{2a}{2b+c}$. While on the left you recover $\sum_{\text{cyc}} \frac{b+c}{\sqrt{(a+b)(a+c)}}$.

Thus in full generality

$$\sum_{\text{cyc}} \frac{b+c}{\sqrt{(a+b)(a+c)}} < 2 \sum_{\text{cyc}} \frac{2a}{2b+c},$$

Therefore, the answer is (D) $<$

Irrelevant theorem provided: Theorem Zdravko (Mean inequalities): Let $a_1, a_2, \dots, a_n$ be positive real numbers. The numbers

$$QM = \sqrt{\frac{a_1^2 + a_2^2 + \dots + a_n^2}{n}}, \quad AM = \frac{a_1 + a_2 + \dots + a_n}{n}$$

$$GM = \sqrt[n]{a_1 a_2 \dots a_n} \quad \text{and} \quad HM = \frac{n}{\frac{1}{a_1} + \frac{1}{a_2} + \dots + \frac{1}{a_n}}$$

are called the quadratic, arithmetic, geometric and harmonic mean for the numbers $a_1, a_2, \dots, a_n$, respectively, and we have

$$QM \geq AM \geq GM \geq HM$$

Equalities occur if and only if $a_1 = a_2 = \dots = a_n$.

Response by o4-mini with theorem provided (Incorrect response ✗):

First, when $a = b = c = 1$, we have

$$LHS = \sum_{\text{cyc}} \frac{b+c}{\sqrt{(a+b)(a+c)}} = 3, \quad RHS = 2 \sum_{\text{cyc}} \frac{2a}{2b+c} = 4.$$

Hence at $a = b = c$,

$$LHS = 3 < RHS = 4.$$

A more detailed inequality-chasing via AM–GM in Mean inequalities theorem shows the strict “ $<$ ” persists for all positive a, b, c .

Therefore, the answer is (D) $<$

Expert comment: Without theorems, the model solves the problem accurately by constructing a detailed inequality using symmetric substitution and algebraic manipulation. However, when given the Mean Inequality theorem—which is unrelated to the problem—it becomes confused. Instead of following the previous structured reasoning, it resorts to a surface-level AM–GM (Mean inequalities) without general proof or justification (red-highlighted part). This highlights the importance of supplying relevant theorems.

Retrieving training problems as demonstrations. Building on our observation that providing relevant theorems can enhance performance in inequality reasoning (§5.4, §D.3, §D.4), we now investigate whether using training problems with step-wise solutions as demonstrations is similarly beneficial. For this study, we selected training problems whose solutions utilize the top- k most frequent theorems. As shown by the overall accuracy in Figure 15, Grok 3 mini’s performance improves from a baseline of 10% (with no demonstration problem) to 13% when provided with one such problem. However, its accuracy drops sharply to 3% when two problems are used as demonstrations. Similarly, Gemini 2.5 Pro peaks at 53% accuracy with one demonstration problem, declining to 45% with two. o4-mini reaches 23% accuracy with one demonstration problem, a 3% increase from its 20% baseline (without demonstrations).

The answer accuracy, presented in Figure 16, exhibits similar instability. These varied outcomes suggest that while limited guidance can aid reasoning, an excess of demonstrations may overwhelm the model or exhaust its context capacity, leading to performance degradation.

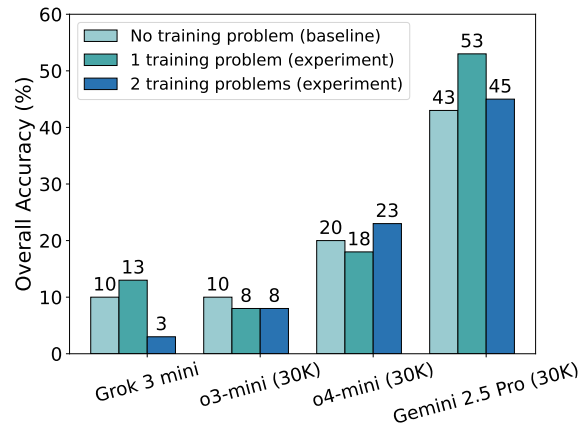


Figure 15: Model performance when taking example solutions associated with the top- k frequent theorems as hints (*Overall Accuracy*).

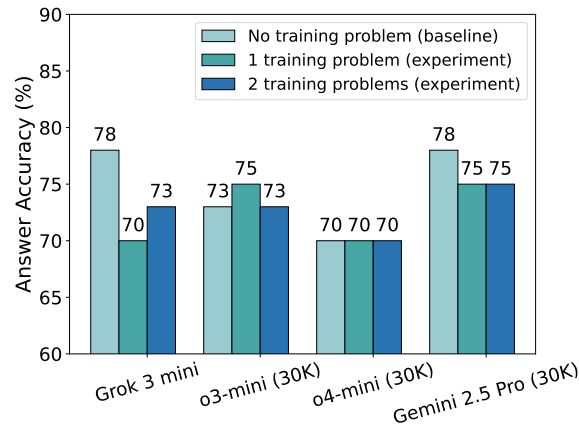


Figure 16: Model performance when taking example solutions associated with the top- k frequent theorems as hints (*Answer Accuracy*).

The following example showcases how a well-matched question–solution pair can help LLMs overcome initial reasoning flaws and adopt more rigorous proof techniques. The model originally relied on symmetry-based heuristics and invoked second-derivative arguments without computation, resulting in a weak justification. After being presented with a relevant example involving convexity and Jensen’s inequality, it successfully reproduced the key steps—verifying convexity analytically and

applying the correct inequalities with full justification. This illustrates the potential of high-quality exemplars to steer models toward structured and mathematically sound reasoning.

Retrieving Training Problems as Demonstrations: Example

Question: Let $a, b > 0$ such that $a + b = 1$. Determine the largest constant C such that the following inequality holds for all a, b :

$$\left(a + \frac{1}{a}\right)^2 + \left(b + \frac{1}{b}\right)^2 \geq C.$$

Ground truth: $C = \frac{25}{2}$

Response by o4-mini without examples (Incorrect response ✗):

Define $f(a) = \left(a + \frac{1}{a}\right)^2 + \left((1-a) + \frac{1}{1-a}\right)^2$, $0 < a < 1$.

1. The function f is symmetric about $a = \frac{1}{2}$. Hence any interior extremum must occur at $a = \frac{1}{2}$.
2. At $a = \frac{1}{2}$, we have $f\left(\frac{1}{2}\right) = 2 \cdot \frac{25}{4} = \frac{25}{2}$.
3. Checking the second derivative shows this is indeed a minimum. The answer is $C = \frac{25}{2}$.

Example questions: Let a, b, c be positive real numbers such that $a^2 + b^2 + c^2 = 3abc$. Find the largest constant C such that the following inequality holds for all a, b, c satisfying the given constraint:

$$\frac{a}{b^2c^2} + \frac{b}{c^2a^2} + \frac{c}{a^2b^2} \geq \frac{C}{a+b+c}$$

Example solution: The given inequality is equivalent to

$$(a^3 + b^3 + c^3)(a + b + c) \geq 9a^2b^2c^2$$

Applying the Cauchy-Schwarz inequality we have

$$(a^3 + b^3 + c^3)(a + b + c) \geq (a^2 + b^2 + c^2)^2$$

Since $a^2 + b^2 + c^2 = 3abc$ we obtain

$$(a^3 + b^3 + c^3)(a + b + c) \geq (a^2 + b^2 + c^2)^2 = (3abc)^2 = 9a^2b^2c^2$$

Equality holds if and only if $a = b = c = 1$. Therefore, the answer is $C = 9$.

Response by o4-mini with examples (Correct response ✓): Let $a, b > 0$ with $a + b = 1$. Define the function $f(a) = \left(a + \frac{1}{a}\right)^2 + \left(b + \frac{1}{b}\right)^2$, $b = 1 - a$.

Convexity argument: consider the function $h(x) = \left(x + \frac{1}{x}\right)^2$ ($x > 0$).

Compute its second derivative: $h(x) = x^2 + 2 + \frac{1}{x^2}$, $h''(x) = 2 + 6x^{-4} > 0$, $\forall x > 0$. Hence h is convex on $(0, \infty)$. Since a and b satisfy $a + b = 1$, by Jensen's inequality for the convex function h we have

$$\frac{h(a) + h(b)}{2} \geq h\left(\frac{a+b}{2}\right) = h\left(\frac{1}{2}\right) = \left(\frac{1}{2} + 2\right)^2 = \left(\frac{5}{2}\right)^2 = \frac{25}{4}.$$

Therefore $f(a) = h(a) + h(b) \geq 2 \cdot \frac{25}{4} = \frac{25}{2}$, i.e. The answer is $C = \frac{25}{2}$.

Expert comment: The initial solution lacks rigor: it incorrectly assumes symmetry implies a unique minimum, and references the second derivative without computing it (red-highlighted part). However, after being provided with a related example that uses a careful convexity argument and the Cauchy-Schwarz inequality, the model is able to imitate the correct reasoning. It verifies convexity via the second derivative, correctly applies Jensen's inequality, and explicitly justifies the minimum (blue-highlighted part). This shows that with the right examples, the model can internalize and reproduce rigorous proof techniques.

D.5 Self-improvement via Critic as Feedback

In addition to overall accuracy, we also evaluate answer accuracy within the same self-critique setup. Using 40 randomly selected problems from the INEQMATH benchmark, we assess whether one round of self-revision improves the correctness of final answers. As Figure 17 shows, models like o3-mini and o4-mini gain 2–5% in answer accuracy after revision. This result further supports self-critique as a lightweight and supervision-free approach to improve solution reliability in inequality problems.

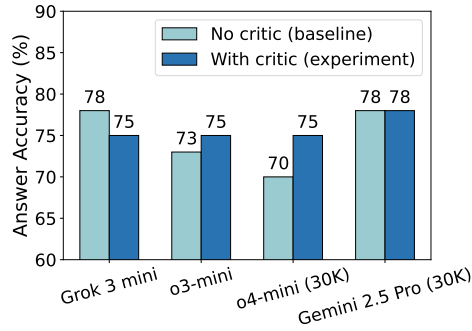


Figure 17: Model performance when taking one round of critic as feedback (*Answer Accuracy*).

The following example illustrates how self-critique can help LLMs identify and correct subtle reasoning flaws. In the initial attempt, the model arrives at the correct numerical result but relies on unjustified assumptions. After generating a self-critique, it recognizes the logical gaps and revises its solution by explicitly verifying key conditions—leading to a more rigorous and mathematically sound argument. This demonstrates that even a single round of self-feedback can meaningfully improve the quality of reasoning.

Self-improvement via Critic as Feedback: Example

Problem: Let $a, b > 0$ such that $a + b = 1$. Determine the largest constant C such that the following inequality holds for all a, b :

$$\left(a + \frac{1}{a}\right)^2 + \left(b + \frac{1}{b}\right)^2 \geq C.$$

Ground truth: $C = \frac{25}{2}$

Initial response by o4-mini (Incorrect response ✗):

Define $f(a) = \left(a + \frac{1}{a}\right)^2 + \left((1-a) + \frac{1}{1-a}\right)^2$, $0 < a < 1$.

1. The function f is symmetric about $a = \frac{1}{2}$. Hence any interior extremum must occur at $a = \frac{1}{2}$.
2. At $a = \frac{1}{2}$, we have $f\left(\frac{1}{2}\right) = 2 \cdot \frac{25}{4} = \frac{25}{2}$.

3. Checking the second derivative shows this is indeed a minimum. The answer is $C = \frac{25}{2}$.

(Continued)

Self-improvement via Critic as Feedback: Example (Continued)**Self-critic feedback by o4-mini:**

These are what make this solution incompletely justified:

- symmetry alone guarantees $f'(1/2) = 0$, but it does *not* by itself rule out other critical points $a \neq 1/2$. In general, a symmetric C^1 -function can have multiple stationary points symmetrically placed about the center.
- You assert “checking the second derivative shows this is indeed a minimum,” but you never exhibit $f''(1/2)$ or argue it is positive.

Refined solution by o4-mini (Correct response ✓):

Define the function $f(a) = \left(a + \frac{1}{a}\right)^2 + \left((1-a) + \frac{1}{1-a}\right)^2$, for $a \in (0,1)$. Let $g(x) = \left(x + \frac{1}{x}\right)^2$ for $x > 0$. Then $f(a) = g(a) + g(1-a)$.

Compute:

$$g(x) = x^2 + 2 + \frac{1}{x^2}, \quad g'(x) = 2x - \frac{2}{x^3}, \quad g''(x) = 2 + \frac{6}{x^4} > 0.$$

Hence, g is strictly convex on $(0, \infty)$.

Since $f''(a) = g''(a) + g''(1-a)$ and each term is positive for $a \in (0,1)$, we have $f''(a) > 0$. Thus, f is strictly convex on $(0,1)$. A strictly convex C^1 -function on an open interval has at most one critical point, which, if it exists, is its global minimum. We compute:

$$f'(a) = g'(a) - g'(1-a) = 0$$

By symmetry, we can solve it and get $a = \frac{1}{2}$. Hence the unique solution of $f'(a) = 0$ in $(0,1)$ is $a = \frac{1}{2}$.

Therefore, the smallest possible value is $C = f\left(\frac{1}{2}\right) = \frac{25}{2}$.

Expert comment: The initial solution gets the correct value $C = \frac{25}{2}$, but lacks rigor: it incorrectly assumes symmetry implies a unique minimum, and references the second derivative without computing it (red-highlighted part). The critic accurately points out these gaps. The refined solution corrects them by proving convexity through second derivatives and confirming $a = \frac{1}{2}$ is the unique minimizer through first derivatives (blue-highlighted part). This demonstrates that self-criticism effectively improves the mathematical soundness of the argument.

E Limitations

While our work introduces a novel dataset and evaluation judges for LLM-based inequality proving, we acknowledge several limitations that warrant discussion and offer avenues for future research.

Potential for data contamination. Although we took significant measures to mitigate data leakage by commissioning novel test problems curated by experts, keeping ground truth answers private, and utilizing an online leaderboard for evaluation, a residual risk of contamination remains. LLMs possess vast training corpora, and it is possible they have encountered problems with similar structures or underlying principles during pre-training, potentially inflating performance beyond true generalization capabilities. Our expert curation and review process aimed to minimize this, but perfect isolation from prior knowledge is challenging to guarantee.

Training dataset scale and scope. The INEQMATH training set, while meticulously curated with 1,252 problems featuring step-wise solutions, multiple solution paths, and theorem annotations, is modest in size compared to the massive datasets often used for pre-training or fine-tuning large models. We prioritized quality and depth (step-wise solutions, theorems) to the challenging Olympiad-level domain over sheer quantity. While sufficient for benchmarking current models, post-training, and exploring test-time techniques, this scale might be insufficient for training highly specialized models from scratch or for capturing the full diversity of inequality types. Future work could focus on scaling up the dataset while maintaining quality, potentially through community contributions.

Inherent inaccuracies in LLM-as-judge evaluation. Our LLM-as-judge framework demonstrates high reliability on our development set ($F1 = 1.0$ for final-answer judge, > 0.9 average for step-wise judges). However, while significantly more scalable than human expert evaluation, these judges are still imperfect. As illustrated by examples in §C.8, they can occasionally misinterpret complex reasoning, overlook subtle logical flaws, or fail to correctly assess nuanced mathematical arguments. The current set of step-wise judges targets common failure modes but does not cover all possible error types, such as the correctness of complex symbolic transformations or the optimal choice of strategy. Potential improvements include using more powerful (but potentially more expensive) LLMs as judge backends (e.g., o3), developing specialized judges trained on annotated errors, or adding judges for specific mathematical operations like symbolic manipulation verification.

Mitigation, not elimination, of answer guessability. The inclusion of step-wise judges significantly mitigates the issue of models guessing the correct final answer without sound reasoning. However, it does not eliminate this possibility entirely. A model might still arrive at the correct bound or relation through chance or heuristics and support it with plausible-sounding, yet flawed, intermediate steps capable of misleading one or more judges. The requirement to pass all judges reduces this risk, but the fundamental challenge of distinguishing genuine mathematical insight from convincing yet spurious reasoning remains.

Computational cost of evaluation. While more efficient than manual expert grading, our multi-judge evaluation protocol is computationally more intensive than simple final-answer checking (e.g., string matching). Evaluating each solution requires multiple LLM inferences (one for the final answer, four for step-wise checks). This cost scales linearly with the number of models and problems being evaluated and could become a factor in very large-scale benchmarking efforts.

F Broader Impacts

This research focuses on advancing the mathematical reasoning capabilities of LLMs, specifically in the domain of inequality proving. While the work is primarily foundational and unlikely to lead directly to malicious applications such as disinformation or surveillance, potential negative societal impacts could arise from the misuse or misinterpretation of the technology. The most significant risk stems from over-reliance on LLM-generated proofs that may appear correct superficially (achieving high answer accuracy) but contain critical logical flaws, as demonstrated by the sharp drop in performance under our step-wise evaluation. If such flawed proofs were uncritically accepted in fields requiring mathematical rigor, such as scientific modeling, engineering design, or financial analysis, it

could lead to incorrect conclusions, faulty systems, or economic miscalculations. Our contribution of a rigorous, step-wise evaluation methodology serves as a potential mitigation strategy by promoting transparency and enabling the identification of fragile reasoning chains, thereby encouraging cautious deployment and emphasizing the need for verification, especially in high-stakes applications. The public release of the `INEQMATH` benchmark further supports community efforts in understanding and improving the reliability of LLM reasoning.