

SCIAGENT: Tool-augmented Language Models for Scientific Reasoning

Anonymous ACL submission

Abstract

Scientific reasoning poses an excessive challenge for even the most advanced Large Language Models (LLMs). To make this task more practical and solvable for LLMs, we introduce a new task setting named *tool-augmented scientific reasoning*. This setting supplements LLMs with scalable toolsets, and shifts the focus from pursuing an omniscient problem solver to a proficient tool-user. To facilitate the research of such setting, we construct a tool-augmented training corpus named MATH-FUNC which encompasses over 30,000 samples and roughly 6,000 tools. Building on MATH-FUNC, we develop SCIAGENT to retrieve, understand and, if necessary, use tools for scientific problem solving. Additionally, we craft a benchmark, SciTOOLBENCH, spanning five scientific domains to evaluate LLMs’ abilities with tool assistance. Extensive experiments on SciTOOLBENCH confirm the effectiveness of SCIAGENT. Notably, SCIAGENT-MISTRAL-7B surpasses other LLMs with the same size by more than 13% in absolute accuracy. Furthermore, SCIAGENT-DEEPMATH-7B shows much superior performance than ChatGPT.

1 Introduction

Scientific reasoning (Ouyang et al., 2023; Zhao et al., 2023) aims to comprehend and make decisions regarding problems among STEM (*Science, Technology, Engineering and Mathematics*) domains. It is a fundamental aspect of intelligence, a demanding capability of Large Language Models (LLMs), and a notoriously challenging task. For instance, even GPT-4 (OpenAI, 2023) achieves only 50% and 35% accuracy on TheoremQA (Chen et al., 2023b) and SciBench (Wang et al., 2023b), respectively. Regarding open-source LLMs such as LLaMA-2 (Touvron et al., 2023) and CodeLlama (Rozière et al., 2023), their performances are only about 10% accuracy or even less.

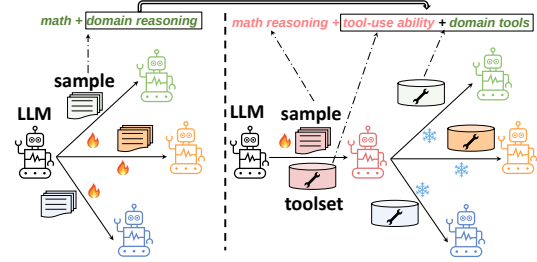


Figure 1: Two paradigms for scientific reasoning. Different colors represent different scientific domains. **Left:** Collecting annotations and fine-tuning LLMs domain by domain. **Right:** Our proposed *tool-augmented* setting. LLMs are fine-tuned on math-related, tool-augmented samples (color in red). When adapting LLMs to a specific domain, a pluggable and domain-specific toolset is attached. No additional fine-tuning is further required.

The challenge in scientific reasoning arises from the need for both mathematical (math) and domain-specific reasoning abilities. To address the physical problem in Figure 3, for example, it is necessary to both understand *Malus’ law* (domain knowledge) for analyzing the intensity of polarized light, and possess quantitative ability for calculating the light intensity ratios. A natural approach involves collecting annotations and fine-tuning LLMs to enhance their math and domain-specific reasoning abilities, as depicted in Figure 1 (left). However, annotating scientific reasoning problems is extremely expensive. What is worse, adapting LLMs to a new domain demands a fresh round of annotation and fine-tuning, rendering this approach impractical.

In this paper, we draw inspirations from *tool learning* (Qin et al., 2023a) to enhance LLMs’ scientific reasoning capabilities. Instead of solving scientific problem from scratch, humans have summarized and wrapped various points as generalized and well-documented functions in scientific computing softwares, such as Matlab, WolframAlpha, SymPy, etc. These functions¹, which could be

¹In this work, tools refer to Python functions. We use tools and functions interchangeably unless otherwise specified.

equivalently viewed as external tools, greatly facilitate math-adept users to solve difficult scientific problems. In analogy with humans, we do not pursue an omniscient **solver** across various scientific domains. Instead, we assume the access to domain-specific toolsets and pursue a unified, generalized LLM-based **tool-user** as shown in the Figure 1 (right). This approach tackles domain-specific reasoning challenges by enabling LLMs learn to use a reusable and scalable toolkit. It alleviates the reasoning challenges of LLMs by concentrating solely on enhancing their tool-use abilities. These abilities are not only easier to acquire but also applicable across a variety of scientific fields. By attaching domain-specific toolsets, our tool-users can be readily adapted to different fields without the need for additional in-domain fine-tuning.

This work focuses on developing and benchmarking the ability of LLMs in scientific reasoning **with the help of tools**. We envision a scenario where LLMs have access to a domain-specific toolset, comprising various specialized functions. Upon this scenario, we propose a complete framework of dataset construction, model training and evaluation. Given a scientific question, LLMs are supposed to retrieve functions from the toolset and optionally incorporate functions into the formulated solution. We employ an automatic pipeline featuring GPT-4 to compile a large-scale, math-related, tool-augmented training corpus named as MATHFUNC. This corpus is designed to enable LLMs to learn both essential math skills and how to retrieve, understand and use functions properly. As a result, MATHFUNC contains 31,375 samples and equipped with a toolset encompassing 5,981 generalized and well-documented functions. We detail this training corpus in Section 3.

We fine-tune open-source LLMs on MATHFUNC to develop tool-augmented agents named SCIAAGENT detailed in Section 4. As shown in Figure 3, SCIAAGENT firstly generate a *high-level planning* in response to a given question. The agents then use this plan, along with the question, to retrieve functions from the given toolset. Leveraging these retrieved functions, the agents further complete the *low-level action* integrating natural language and Python code. Finally the agents execute the code to complete the problem at hand.

To benchmark the tool-use abilities in scientific reasoning, we develop a new benchmark named SCIToolBENCH as described in Section 5. Building upon TheoremQA (Chen et al., 2023b) and

SciBench (Wang et al., 2023b), it has 856 questions covering five domains: *Mathematics*, *Physical*, *Chemistry*, *EECS*, and *Finance*. It also contains five domain-specific toolsets comprising a total of 2,446 functions. We evaluate SCIAAGENT on SCIToolBENCH and another benchmark derived from CREATOR-challenge (Qian et al., 2023). Experimental results demonstrate that our agents present remarkable scientific reasoning capabilities. Notably, SCIAAGENT-MISTRAL-7B surpasses the best comparable open-source LLMs by an absolute 13.4% accuracy, and SCIAAGENT-DEEPMATH-7B outperforms ChatGPT by a large margin. We also conduct an extensive analysis of the benefits and limitations of SCIAAGENT series, providing valuable insights for future research.

2 Preliminary

Related Work. Current methods (Chen et al., 2023b; Xu et al., 2023b; Ouyang et al., 2023), especially those based on open-source LLMs, perform far from satisfactory on scientific reasoning benchmarks (Chen et al., 2023b; Wang et al., 2023b). We attribute it to the scarcity of annotated samples across diverse scientific domains. As a comparison, LLMs present much more remarkable performance on math problems (Yue et al., 2023b; Gou et al., 2023b; Azerbayev et al., 2023) due to the abundant training corpora and/or annotations. Different from concurrent work (Zhang et al., 2024) which collects physics and chemistry annotations, we do not pursue a problem-solver on some specific scientific domains. Instead, we consider to develop a generalized tool-user being proficient on solving diverse scientific problems with the aid of tools. Following previous work on math domain (Qian et al., 2023; Cai et al., 2023; Yuan et al., 2023a), the tools here refer to Python functions. Please see more detailed literature review in Appendix A.

Task Formulation. Given a scientific domain D (e.g., physics), *tool-augmented* scientific reasoning task assumes access to (1) a question $q \in D$ and (2) a toolset F_D . F_D encompasses large amounts of well-documented, domain-specific functions $\{f_1, \dots, f_m\}$. Our objective is to develop an agent $\mathcal{M}$ which selectively use functions in F_D to enhance the answering for the question q .

3 Training Corpus: MATHFUNC

To our best knowledge, there are no readily available tool-augmented datasets in scientific reason-

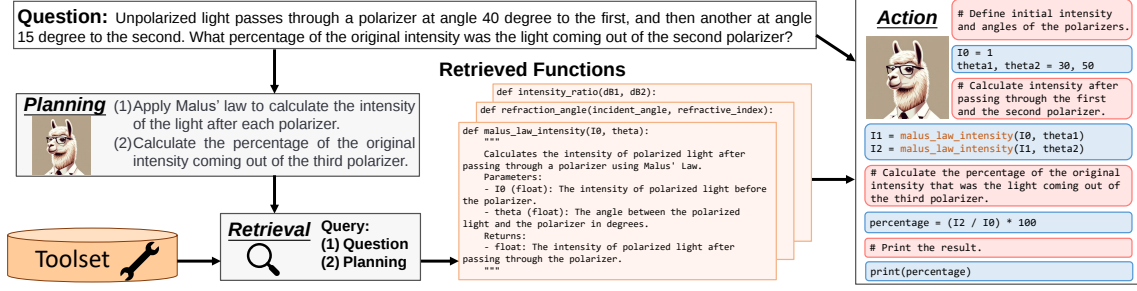


Figure 3: The model architecture of SCIAGENT. Given a domain-specific toolset , our agent answers the question through four consecutive modules. (1) **Planning** : provides a high-level plan for this problem. (2) **Retrieval** : retrieves related functions from attached toolset. (3) **Action** : generates a low-level solution interleaving rationale and program. The program uses the retrieved functions if necessary. (4) **Execution** : calls Python executor to run the program and outputs the final answer. Not included in this figure for simplicity.

functions, all of which are derived from other questions, in the right side of Figure 2.

Retriever. The cross-retrieval strategy necessitates a retriever because it is impractical to enumerate thousands of functions in $F \setminus \tilde{F}_q$. We train a dense retriever R (③ in Figure 2). We concatenate the question q and the generated planning G_q as the query, and view the generated functions $\tilde{F}_q$ as the keys. See details about R in Appendix B.1.

Solution Generation. Upon the toolset F and the retriever R , we retrieve three functions as F_q :

$$F_q = R([q, G_q]; F \setminus \tilde{F}_q)$$

Then we employ GPT-4 to write solutions which optionally call functions in F_q to generate the solution S_q (④). The prompt used is illustrated in Appendix F.3. We explicitly point out in the prompt that $f \in F_q$ should be called if and only if when they do lower the difficulty of problem solving. It mitigates the over-exploitation of function calling in S_q and increases the robustness of models fine-tuned on these samples. Specifically, we firstly use GPT-4 with greedy decoding to generate solutions. For those failing to yield correct answers, we further apply nucleus sampling (Holtzman et al., 2020) with 5 repeat times and 0.6 temperature. We filter wrong solutions and collect remaining 6,229 samples as our function-augmented solutions.

In parallel, we use GPT-4 to generate function-free solutions. Though not indispensable, we expect them to further enhance the math reasoning, and accordingly the scientific reasoning, abilities of LLMs. We collect a total of 24,946 function-free solutions nucleus sampling with 5 repeat times and 0.6 temperature. These samples share similar format as ToRA-corpus (Gou et al., 2023b), and do not retrieve/use any functions, i.e., $F_q = \emptyset$.

4 Model: SCIAGENT

We develop SCIAGENT for tool-augmented scientific reasoning task. It could make plan, retrieve functions, and leverage retrieved functions to facilitate the reasoning. We describe its inference procedure and training approach as below.

4.1 Overview

As shown in Figure 3, SCIAGENT comprises four successive modules.

Planning. This module provides a high-level profile for each question: $G_q = \mathcal{M}_{\text{planning}}(q)$. Such planning instructs a more targeted retrieval process.

Retrieval. Given the question and generated planning G_q , the retriever $\mathcal{M}_{\text{retrieval}}$ is introduced to retrieve related functions from the domain-specific toolset: $F_q = \mathcal{M}_{\text{retrieval}}([q, G_q]; F_D) \subseteq F_D$.

Action. This module aims to generate low-level solutions. Specifically, the agent produces $S_q = \mathcal{M}_{\text{action}}(q; F_q)$. The solution S_q is interleaved with natural language rationale E_q and program snippet P_q . The program P_q call retrieved functions with proper arguments if necessary.

Execution. This module is simply a Python Executor to run the program P_q for the final answer: $a_q = \text{Python-Executor}(P_q)$.

4.2 Training

Language models are used in three out of four modules in SCIAGENT: planning, retrieval and action. Regarding retrieval, we directly use the retriever R fine-tuned in Section 3.2 as $\mathcal{M}_{\text{retrieval}}$. For planning and action modules, they share the same LLMs: $\mathcal{M} = \mathcal{M}_{\text{planning}} = \mathcal{M}_{\text{action}}$. We fine-tune $\mathcal{M}$ with different instructions to make it act as planning and action modules, respectively. We construct instructions from $d = (q, G_q, F_q, S_q, a_q)$ in MATHFUNC.

$$D_{\text{planning}} = \{(I_{\text{plan}}(q), G_q) | d \in D\}$$

$$D_{\text{action}} = \{(I_{\text{action}}(q, F_q), S_q) | d \in D\}$$

Here I_{plan} and I_{action} are instruction templates for planning and action modules. We show these instructions in Appendix B.2, and mix up them as the training set $D = (D_{\text{planning}} \cup D_{\text{action}})$. Then we apply imitation learning on D to fine-tune $\mathcal{M}$.

$$L_{\mathcal{M}} = \sum_{(X,Y) \in D} -\log \mathcal{P}(Y|X)$$

Implementation We detail the training process of (1) the retriever $\mathcal{M}_{\text{retrieval}}$ and (2) the planner and actor $\mathcal{M}$ in Appendix B.1 and B.2, respectively.

5 Benchmark: SciToolBENCH

There currently exists no benchmark assessing the scientific reasoning capabilities of LLMs **when aided by tools**. To address this gap, we develop a benchmark called SciToolBENCH. Our benchmark covers five domains: *Mathematics (math)*⁵, *Physics*, *Chemistry*, *Finance*, *Electrical Engineering and Computer Science (EECS)*. Each domain is composed of a set of questions and a domain-specific toolset. The toolset consists of abundant generalized, high-quality and well-documented functions. We expect LLMs to retrieve, understand and, if necessary, use functions in it for reasoning.

Table 1: The statistics of our benchmark. **#Func**: Number of functions. **#Pos./ #Neg.**: The number of positive/negative functions in the toolset. **FPQ** (function per question): The number of derived positive functions from each question.

	# Question	# Func	# Pos. / # Neg.	Avg. FPQ
Math	434	1072	511 / 561	1.47
Physics	156	534	243 / 291	1.63
Chemistry	118	366	155 / 211	1.34
Finance	66	253	97 / 156	1.62
EECS	82	221	97 / 124	1.68
All	856	2446	1103 / 1343	1.51

5.1 Dataset Overview.

The statistics of SciToolBENCH are presented in Table 1. It comprises a total of 856 questions and 2,446 functions spanning across 5 scientific domains. Notably, SciToolBENCH differs from previous tool-based benchmarks, such as Creation

⁵Our benchmark contains college-level questions on calculus, differential equations, group theory, etc, which are different from the questions in our training corpus MATHFUNC.

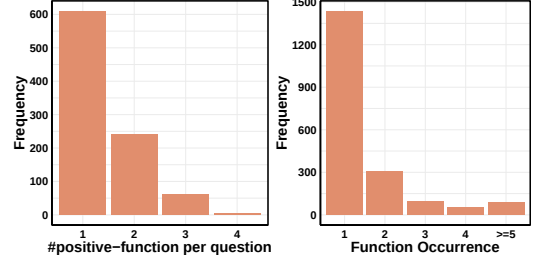


Figure 4: **Left**: Histogram of FPQ (function per question). Higher values indicate greater composability. **Right**: Histogram of function occurrence. Higher values indicate more generalization and wider application.

Challenge (Qian et al., 2023), in several aspects: (1) Our benchmark encompasses a diverse range of scientific domains. (2) The tools provided are both composable and generalized across different questions. As indicated in Table 1, each question requires an average of 1.51 functions for resolution. And as shown in Figure 4, over 500 functions are designed to be applicable to two or more questions, such as `integrate_function` in math domain, `coulombs_law` in physical domain, and `calculate_pressure_van_der_waals` in chemistry domain. It signifies that the functions in our toolset are not ad-hoc solutions tailored for specific questions. Instead, the effective utilization of the toolset demands significant reasoning abilities of tool-augmented LLMs. Thus we claim this benchmark challenging and practical.

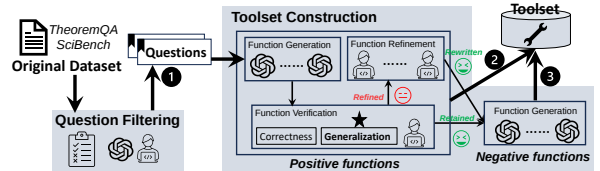


Figure 5: Semi-automatic annotation pipeline for SciToolBENCH. : GPT-4. : Human annotator.

5.2 Dataset Annotation

We design a pipeline shown in Figure 5 to annotate the benchmark. It employs both GPT-4 and human annotators to combine their merits. We introduce it briefly as below and leave details in Appendix D.

Question Filtering: We curate questions from TheoremQA (Chen et al., 2023b) and SciBench (Wang et al., 2023b) to collect 856 questions (① in Figure 5, the same below) in our benchmark.

Toolset Construction: We construct domain-specific toolsets via two cascade modules: positive and negative function construction. We define positive functions (②) as functions directly deriving from questions. The candidate positive functions

Table 2: Main results on two benchmarks. We highlight our SCiAGENT series in **blue**. The best results (among all open-source LLMs, the same below) are in bold face and the second best are underlined.

Model	Size	Toolset	CREATION	SciTOOLBENCH					
				Math	Physics	Chemistry	Finance	EECS	All
ChatGPT	-	✗ ✓	54.6 59.8	33.4 32.0	19.2 31.4	18.6 33.9	53.0 53.0	25.6 48.8	29.6 35.4
GPT-4	-	✗ ✓	60.0 69.8	52.8 63.1	42.9 63.5	47.5 63.6	65.2 80.3	35.4 80.5	49.5 66.2
LLaMA2	7B	✓	12.6	4.3	10.9	8.4	13.6	11.0	8.3
CodeLlama	7B	✗	17.7	6.5	0.6	5.1	4.9	7.6	5.1
CodeLlama	7B	✓	26.1	9.2	8.3	10.2	24.2	25.6	11.9
Llemma	7B	✗	26.4	10.4	4.5	8.5	10.6	7.3	8.8
Llemma	7B	✓	34.3	16.4	21.2	14.4	36.4	22.0	19.1
Mistral	7B	✗	30.1	11.3	4.5	7.6	16.7	6.1	9.5
Mistral	7B	✓	27.6	13.1	13.5	14.4	34.8	19.5	15.6
Deepseek-Coder	7B	✗	36.8	20.3	8.3	5.9	22.7	12.2	15.5
Deepseek-Coder	7B	✓	31.3	21.0	15.4	10.2	30.3	36.6	20.7
Deepseek-Math	7B	✗	44.7	26.5	19.2	17.8	27.3	20.7	23.5
Deepseek-Math	7B	✓	41.3	24.2	24.4	25.4	43.9	42.7	27.7
ToRA-Coder	7B	✗	29.7	26.3	4.5	6.8	9.1	24.4	18.1
ToRA-Coder	7B	✓	21.4	21.7	4.5	5.1	13.6	15.9	15.1
MAMmoTH-Coder	7B	✓	21.6	14.8	18.5	11.0	25.8	40.0	19.7
SCiAGENT-CODER	7B	✓	53.0	30.0	28.3	24.6	39.3	<u>57.3</u>	32.2
SCiAGENT-MISTRAL	7B	✓	54.0	31.3	28.8	22.9	51.5	61.0	34.1
SCiAGENT-DEEPMATH	7B	✓	60.4	41.2	54.5	44.9	57.5	51.2	46.3
LLaMA2	13B	✓	23.3	12.2	11.5	6.8	22.7	14.6	12.4
CodeLlama	13B	✗	23.0	9.9	3.2	1.7	9.1	6.1	7.1
CodeLlama	13B	✓	38.9	12.7	14.7	7.6	33.3	34.1	16.0
ToRA-Coder	13B	✗	30.9	28.6	3.8	4.2	16.7	30.5	20.0
ToRA-Coder	13B	✓	28.0	32.0	2.6	11.9	24.2	35.4	23.6
MAMmoTH-Coder	13B	✓	34.7	21.4	18.6	11.0	25.8	39.0	21.5
SCiAGENT-CODER	13B	✓	<u>54.4</u>	<u>35.0</u>	<u>32.1</u>	<u>28.8</u>	42.4	51.2	<u>35.7</u>

(2) are firstly generated from GPT-4. Then human annotators carefully check them and rewrite and/or remove the unqualified ones. We further automatically construct negative functions (3) based on positive functions to reduce the shortcuts in our benchmark. We finally combine both positive and negative functions as the toolset in our benchmark.

6 Experiments

6.1 Setup

We conduct experiments on SciTOOLBENCH to evaluate the tool-augmented scientific reasoning abilities of LLMs. We also employ CREATION Challenge (Qian et al., 2023) as the second benchmark. It comprises a total of 2,047 samples, with each sample consisting of a question and a ground-truth function. We aggregate all functions to assemble the toolset (thus including 2,047 functions). We report accuracy as the metric in all experiments.

6.2 Baselines

We compare SCiAGENT series with eight open-source LLMs: (1) LLaMA-2 (Touvron et al., 2023), (2) CodeLlama (Rozière et al., 2023), (3) Mis-

tral (Jiang et al., 2023), (4) Llemma (Azerbayev et al., 2023), (5) Deepseek-Coder (Guo et al., 2024), (6) Deepseek-Math (Shao et al., 2024), (7) MAMmoTH-Coder (Yue et al., 2023b) and (8) ToRA-Coder (Gou et al., 2023b). We also list the performance of ChatGPT and GPT-4 for reference. We provide all LLMs the same retriever in Section 3.2 to retrieve functions from toolset (if attached). Please see more details in Appendix C.

6.3 Main Results

We fine-tune CodeLlama, Mistral and Deepseek-Math for yielding SCiAGENT-CODER, SCiAGENT-MISTRAL and SCiAGENT-DEEPMATH. We show their results in Table 2 and observe: (1) Almost all LLMs present improved performance, *i.e.*, 5.3% absolute and 61.6% relative accuracy increase on average, when supplemented with toolsets. It validates the promise of the tool-augmented setting for scientific reasoning. (2) The models fine-tuned on math-related datasets from CodeLlama, *i.e.*, ToRA- and MAMmoTH-Coder, perform better than CodeLlama itself by 5.5% absolute accuracy. It presents the importance of essential math skills among diverse scientific domains. (3)

Table 3: Ablation study on SCIToolBENCH. We report the accuracy of samples across (1) all domains, (2) four domains excluding the math domain (wo. math).

	Planning	Function-augmented solutions	Function-free solutions	Retriever	Accuracy (7B)		Accuracy (13B)	
					All	wo. math	All	wo. math
SciAgent-Coder	✓	✓(cross-retrieval)	✓	✓	32.2	34.6	35.7	36.5
Intermediate variants 1-3	✗	✓(cross-retrieval)	✓	✓	30.3	33.9	32.8	34.4
	✗	✓(direct-use)	✓	✓	17.8	17.3	26.6	31.0
	✗	✗	✓	✓	26.3	26.1	30.4	31.7
CodeLlama wo. retriever	✗	✗	✗	✓	11.9	14.7	16.0	19.4
	✗	✗	✗	✗	5.1	3.8	7.1	4.3

Our agents consistently outperform other open-source LLMs by a large margin. Notably, SCIAgent-Coder surpasses the most competitive baseline, MAMmoth-Coder, by absolute accuracy of 12.5% and 14.2% on the 7B and 13B versions. (4) Our strongest agent, SCIAgent-DEEPMATH-7B, substantially outperforms ChatGPT with toolset (46.3% v.s. 35.4%) and shows comparable results to GPT-4 without toolset (46.3% v.s. 49.5%). However, it still falls significantly behind GPT-4 when both are provided with the same tools. Such gap highlights the challenges of tool-augmented scientific reasoning (and our benchmark).

6.4 Ablation Study

We investigate the effectiveness of components in our training data and agent modules. The specific variants we considered are as follows. (1) We remove the planning module in the agent. (2) We additionally drop the cross-retrieval strategy introduced in Section 3.2. In its place, we construct function-augmented solutions directly from $\tilde{F}_q$ and $\tilde{S}_q$. (3) We further remove all function-augmented solutions from our training data, and only keep the solutions without function callings (function-free solutions). (4) We do not fine-tune agents but merely use CodeLlama as $\mathcal{M}_{\text{action}}$ for inference. (5) We drop the retriever to disable the LLMs’ tool-use abilities. Equivalently, it degrades to the baseline of CodeLlama + PoT (Chen et al., 2023a) prompting.

We illustrate the performance of our agents and their ablated variants in Table 3. We observe that (1) Planning module significantly improves scientific reasoning abilities. As detailed and targeted queries for the retriever, the generated plans increase the relatedness of retrieved functions. For instance, the function’s Recall@3 increases from 48.3% to 53.2% in physics domain, and from 37.3% to 39.8% in chemistry domain. (2) The use of the cross-retrieval strategy is essential. Otherwise, the function-augmented solutions directly from $\tilde{F}_q$ and

$\tilde{S}_q$ degrade the performance because they are too artificial and ad-hoc to teach LLMs using functions properly. (3) The absence of function-augmented solutions results in a performance drop (row 1 v.s. row 4 in Table 3) of 5.9% and 5.3% in absolute accuracy for 7B and 13B LLMs, respectively. It underscores the critical role of function-augmented solutions to enhance LLMs’ tool-use abilities, and the necessity of our MATHFUNC corpus. (4) The removal of function-free solutions (row 4 v.s. row 5) leads to an absolutely 14.4% accuracy decrease. Specifically focusing on non-math samples, there is a notable performance drop of about 12% as well. This clearly demonstrates the fundamental importance of math skills in diverse scientific reasoning tasks, and highlights how our math-related samples enhance LLMs’ capabilities in this area. (5) Performance significantly declines when the retriever is removed. It illustrates that the retrieval module is crucial for accessing the appropriate functions from large-scale toolsets.

6.5 Analysis

Robustness of Toolsets. We acknowledge the construction and maintenance of toolsets is sometime challenging. Therefore, we stress the importance of our agents’ robustness. If a sub-par toolset were provided, an robust agent should at the very least perform comparably, if not better, than other competitive LLMs without tool-use. To evaluate the robustness of SCIAgent-Coder, we simulate two sub-par settings. (1) weak-related: for each question, we restrict the agents from retrieving functions that are directly derived from it. This setting greatly decreases the likelihood of retrieving a proper function from the toolset. (2) unrelated: we completely remove the domain-specific toolset in SCIToolBENCH. As a substitution, we provide the unrelated toolset constructed in MATHFUNC.

We compare our agents with two competitive LLMs, *i.e.*, ToRA-Coder and MAMmoth-Coder,

Table 4: Accuracy on SCIAGENT with sub-par toolsets. **WR**: weak-related toolsets. **UR**: unrelated toolsets. **NA**: No toolset. The subscripts indicate the difference from the best LLMs (wo. toolsets) each column.

Model	Toolset	Accuracy (7B)		Accuracy (13B)	
		All	wo.math	All	wo. math
SciAgent-Coder	WR	18.8 _{+0.7}	18.0 _{+8.3}	24.6 _{+4.6}	19.9 _{+7.6}
	UR	14.7 _{-3.7}	10.7 _{+1.0}	20.3 _{+0.3}	14.7 _{+2.4}
MAmmo-C	NA	12.7	9.0	16.4	12.3
ToRA-C	NA	18.1	9.7	20.0	11.1

in above two settings. As shown in Table 4, (1) SCIAGENT series with unrelated toolsets present comparable performance with the two LLMs. In other words, our tool-augmented agents are unlikely to degrade the performance even under the extreme scenarios. (2) Our agents with weak-related toolsets significantly outperform the two LLMs, which further validates the robustness.

The Effect of Retriever Quality. We explore the effect of retriever quality on the ending performance. We substitute our fine-tuned retriever in SCIAGENT series by two competitive variants: SimCSE (Gao et al., 2021) and Contriever (Izacard et al., 2021). As shown in Figure 6 (top), our retriever surpasses the other two. It shows that fine-tuning on the math domain benefits the retrieval of tools in the generalized scientific domains.

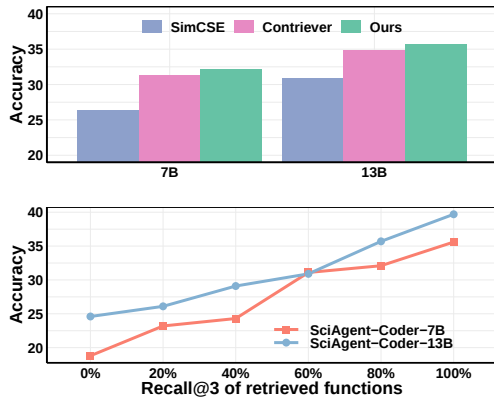


Figure 6: **Top**: Performance of SCIAGENT-CODER on SCITOOLBENCH with different retriever variants. **Bottom**: Relationship between the performance and the hit@3 of retrieved functions (artificially controlled).

We further dive deep into the relationship between the hit ratio of tools and the agents’ performance. To this end, we manually control the hit@3 ratio by artificially adding/removing the positive functions to/from the retrieved list. Results in Figure 6 (bottom) show a clearly positive correlation between the hit ratio and the task accuracy. It illustrates that the retrieved functions facilitate the reasoning of scientific problems. However, we still

observe a limit (40% accuracy) when the hit ratios reaching 100%, showing the challenge of scientific reasoning even when aided by tools. We hope the future work to bridge this performance gap.

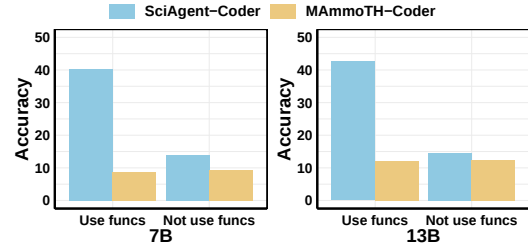


Figure 7: The performance of SCIAGENT-CODER (w. toolset) and MAMMOTH-CODER (wo. toolset) on samples which (1) use and (2) not use retrieved functions.

How the Retrieved Functions Benefit. To assess how the retrieved functions aid in the reasoning process of LLMs, we divided the samples into two subsets based on whether our agents use the retrieved functions to solve the problems. We evaluate the performance of these two subsets respectively, comparing with MAMMOTH-Coder series (without tool-use). The results in Figure 7 reveal a two-fold benefit: (1) For samples where functions are explicitly called to solve the questions, our agents demonstrate a substantial 25% improvement in absolute accuracy over LLMs that do not have access to functions. (2) Even for samples that do not directly use functions in their written program, we still observe a slight improvement. It suggests that our agents are capable of learning from retrieved functions as a reference, and then imitate these functions to write their own programs. For instance, example in Figure 12 shows the agents learn how to use `scipy.integrate` by observing the retrieved function `average_value_of_function(...)`.

7 Conclusion

This work proposes *tool-augmented* scientific reasoning, a task aiming to solve challenging scientific problems aided by generalized and scalable tools. To facilitate and evaluate the scientific tool-use abilities of LLMs, we construct a math-related, tool-augmented training corpus MATHFUNC and a benchmark SCITOOLBENCH covering 5 scientific domains. Additionally, we develop open-source agents, SCIAGENT series, as competitive baselines. Extensive experiments reveal that our agents exhibit tool-use abilities exceeding ChatGPT in scientific reasoning tasks.

Limitations

The primary limitation of our work comes from the way we compile the toolsets in SciToolBench. These tools are constructed directly based on the benchmark’s questions, raising concerns about potential information leakage. To address this, we invest significant human effort in our annotation process as detailed in Appendix D.2. We manually review and, if necessary, revise all derived functions to ensure their generalizability and quality. As shown in Figure 6 (bottom), our agents achieve only about 40% accuracy when we provide each question the exact function from which it derives (*i.e.*, 100% hit ratio). It not only highlights the inherent challenge of scientific reasoning tasks, but also suggests that our benchmark suffers minimal impact from the potential information leakage.

We partly attribute this limitation to the absence of a training corpus among scientific (excluding math) domains. The scarcity of annotated solutions for scientific reasoning problems makes it unfeasible to set aside a portion of questions in our benchmark for tool creation. In future work, we plan to collect diverse and high-quality scientific annotations which enable us to develop a more practical and robust tool-augmented benchmark.

Ethics Statement

We ensure that SCIToolBench was constructed in compliance with the terms of use of all source materials and with full respect for the intellectual property and privacy rights of the original authors of the texts. We also provide details on the characteristics and annotation steps of SCIToolBench in Section 5 and Appendix D. We believe our created datasets do not cause any potential risks.

References

- Zhangir Azerbayev, Hailey Schoelkopf, Keiran Paster, Marco Dos Santos, Stephen McAleer, Albert Q. Jiang, Jia Deng, Stella Biderman, and Sean Welleck. 2023. *Llemma: An open language model for mathematics*.
- Andres M Bran, Sam Cox, Oliver Schilter, Carlo Baldassari, Andrew D White, and Philippe Schwaller. 2023. *Chemcrow: Augmenting large-language models with chemistry tools*.
- Tianle Cai, Xuezhi Wang, Tengyu Ma, Xinyun Chen, and Denny Zhou. 2023. *Large language models as tool makers*.
- Wenhu Chen, Xueguang Ma, Xinyi Wang, and William W. Cohen. 2023a. Program of thoughts prompting: Disentangling computation from reasoning for numerical reasoning tasks. *Transactions on Machine Learning Research*.
- Wenhu Chen, Ming Yin, Max Ku, Pan Lu, Yixin Wan, Xueguang Ma, Jianyu Xu, Xinyi Wang, and Tony Xia. 2023b. *TheoremQA: A theorem-driven question answering dataset*. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 7889–7901, Singapore. Association for Computational Linguistics.
- Zhipeng Chen, Kun Zhou, Beichen Zhang, Zheng Gong, Xin Zhao, and Ji-Rong Wen. 2023c. *ChatCoT: Tool-augmented chain-of-thought reasoning on chat-based large language models*. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 14777–14790, Singapore. Association for Computational Linguistics.
- Ethan Chern, Haoyang Zou, Xuefeng Li, Jiewen Hu, Kehua Feng, Junlong Li, and Pengfei Liu. 2023. Generative ai for math: Abel. <https://github.com/GAIR-NLP/abel>.
- Yin Fang, Xiaozhuan Liang, Ningyu Zhang, Kangwei Liu, Rui Huang, Zhuo Chen, Xiaohui Fan, and Huanjun Chen. 2023. *Mol-instructions: A large-scale biomolecular instruction dataset for large language models*.
- Shen Gao, Zhengliang Shi, Minghang Zhu, Bowen Fang, Xin Xin, Pengjie Ren, Zhumin Chen, Jun Ma, and Zhaochun Ren. 2023. *Confucius: Iterative tool learning from introspection feedback by easy-to-difficult curriculum*.
- Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021. *SimCSE: Simple contrastive learning of sentence embeddings*. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 6894–6910, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Zhibin Gou, Zhihong Shao, Yeyun Gong, Yelong Shen, Yujia Yang, Nan Duan, and Weizhu Chen. 2023a. *Critic: Large language models can self-correct with tool-interactive critiquing*.
- Zhibin Gou, Zhihong Shao, Yeyun Gong, yelong shen, Yujia Yang, Minlie Huang, Nan Duan, and Weizhu Chen. 2023b. *Tora: A tool-integrated reasoning agent for mathematical problem solving*.
- Daya Guo, Qihao Zhu, Dejian Yang, Zhenda Xie, Kai Dong, Wentao Zhang, Guanting Chen, Xiao Bi, Y. Wu, Y. K. Li, Fuli Luo, Yingfei Xiong, and Wenfeng Liang. 2024. *Deepseek-coder: When the large language model meets programming – the rise of code intelligence*.
- Shibo Hao, Tianyang Liu, Zhen Wang, and Zhiting Hu. 2023. *Toolkengpt: Augmenting frozen language models with massive tools via tool embeddings*.

663	Dan Hendrycks, Collin Burns, Steven Basart, Andy	Pan Lu, Baolin Peng, Hao Cheng, Michel Galley, Kai-	720
664	Zou, Mantas Mazeika, Dawn Song, and Jacob Stein-	Wei Chang, Ying Nian Wu, Song-Chun Zhu, and	721
665	hardt. 2021a. Measuring massive multitask language	Jianfeng Gao. 2023. Chameleon: Plug-and-play com-	722
666	understanding . In <i>9th International Conference on</i>	positional reasoning with large language models .	723
667	<i>Learning Representations, ICLR 2021, Virtual Event,</i>		
668	<i>Austria, May 3-7, 2021</i> . OpenReview.net.		
669	Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul	Haipeng Luo, Qingfeng Sun, Can Xu, Pu Zhao, Jian-	724
670	Arora, Steven Basart, Eric Tang, Dawn Song, and	guang Lou, Chongyang Tao, Xiubo Geng, Qingwei	725
671	Jacob Steinhardt. 2021b. Measuring mathematical	Lin, Shifeng Chen, and Dongmei Zhang. 2023. Wiz-	726
672	problem solving with the math dataset. <i>NeurIPS</i> .	ardmath: Empowering mathematical reasoning for	727
		large language models via reinforced evol-instruct .	728
673	Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and	OpenAI. 2023. Gpt-4 technical report .	729
674	Yejin Choi. 2020. The curious case of neural text		
675	degeneration . In <i>8th International Conference on</i>	Siru Ouyang, Zhuosheng Zhang, Bing Yan, Xuan Liu,	730
676	<i>Learning Representations, ICLR 2020, Addis Ababa,</i>	Jiawei Han, and Lianhui Qin. 2023. Structured chem-	731
677	<i>Ethiopia, April 26-30, 2020</i> . OpenReview.net.	istry reasoning with large language models .	732
678	Yuzhen Huang, Yuzhuo Bai, Zhihao Zhu, Junlei	Liangming Pan, Alon Albalak, Xinyi Wang, and	733
679	Zhang, Jinghan Zhang, Tangjun Su, Junteng Liu,	William Wang. 2023. Logic-LM: Empowering large	734
680	Chuanheng Lv, Yikai Zhang, Jiayi Lei, Yao Fu,	language models with symbolic solvers for faithful	735
681	Maosong Sun, and Junxian He. 2023. C-eval: A	logical reasoning . In <i>Findings of the Association</i>	736
682	multi-level multi-discipline chinese evaluation suite	<i>for Computational Linguistics: EMNLP 2023</i> , pages	737
683	for foundation models .	3806–3824, Singapore. Association for Computa-	738
		tional Linguistics.	739
684	Gautier Izacard, Mathilde Caron, Lucas Hosseini, Se-	Shishir G. Patil, Tianjun Zhang, Xin Wang, and	740
685	bastian Riedel, Piotr Bojanowski, Armand Joulin,	Joseph E. Gonzalez. 2023. Gorilla: Large language	741
686	and Edouard Grave. 2021. Unsupervised dense infor-	model connected with massive apis .	742
687	mation retrieval with contrastive learning .		
688	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Men-	Baolin Peng, Michel Galley, Pengcheng He, Hao Cheng,	743
689	sch, Chris Bamford, Devendra Singh Chaplot, Diego	Yujia Xie, Yu Hu, Qiuyuan Huang, Lars Liden, Zhou	744
690	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	Yu, Weizhu Chen, and Jianfeng Gao. 2023. Check	745
691	laume Lample, Lucile Saulnier, L��lio Renard Lavaud,	your facts and try again: Improving large language	746
692	Marie-Anne Lachaux, Pierre Stock, Teven Le Scao,	models with external knowledge and automated feed-	747
693	Thibaut Lavril, Thomas Wang, Timoth��e Lacroix,	back .	748
694	and William El Sayed. 2023. Mistral 7b .		
695	Qiao Jin, Yifan Yang, Qingyu Chen, and Zhiyong Lu.	Cheng Qian, Chi Han, Yi Fung, Yujia Qin, Zhiyuan	749
696	2023. Genegpt: Augmenting large language models	Liu, and Heng Ji. 2023. CREATOR: Tool creation	750
697	with domain tools for improved access to biomedical	for disentangling abstract and concrete reasoning of	751
698	information .	large language models . In <i>Findings of the Associa-</i>	752
		<i>tion for Computational Linguistics: EMNLP 2023</i> ,	753
699	Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick	pages 6922–6939, Singapore. Association for Com-	754
700	Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and	putational Linguistics.	755
701	Wen-tau Yih. 2020. Dense passage retrieval for open-		
702	domain question answering . In <i>Proceedings of the</i>	Yujia Qin, Shengding Hu, Yankai Lin, Weize Chen,	756
703	<i>2020 Conference on Empirical Methods in Natural</i>	Ning Ding, Ganqu Cui, Zheni Zeng, Yufei Huang,	757
704	<i>Language Processing (EMNLP)</i> , pages 6769–6781,	Chaojun Xiao, Chi Han, Yi Ren Fung, Yusheng Su,	758
705	Online. Association for Computational Linguistics.	Huadong Wang, Cheng Qian, Runchu Tian, Kunlun	759
		Zhu, Shihao Liang, Xingyu Shen, Bokai Xu, Zhen	760
706	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	Zhang, Yining Ye, Bowen Li, Ziwei Tang, Jing Yi,	761
707	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	Yuzhang Zhu, Zhenning Dai, Lan Yan, Xin Cong,	762
708	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	Yaxi Lu, Weilin Zhao, Yuxiang Huang, Junxi Yan,	763
709	Roberta: A robustly optimized bert pretraining ap-	Xu Han, Xian Sun, Dahai Li, Jason Phang, Cheng	764
710	proach .	Yang, Tongshuang Wu, Heng Ji, Zhiyuan Liu, and	765
		Maosong Sun. 2023a. Tool learning with foundation	766
711	Yuliang Liu, Xiangru Tang, Zefan Cai, Junjie Lu,	models .	767
712	Yichi Zhang, Yanjun Shao, Zexuan Deng, Helan Hu,	Yujia Qin, Shihao Liang, Yining Ye, Kunlun Zhu, Lan	768
713	Zengxian Yang, Kaikai An, Ruijun Huang, Shuzheng	Yan, Yaxi Lu, Yankai Lin, Xin Cong, Xiangru Tang,	769
714	Si, Sheng Chen, Haozhe Zhao, Zhengliang Li, Liang	Bill Qian, Sihan Zhao, Lauren Hong, Runchu Tian,	770
715	Chen, Yiming Zong, Yan Wang, Tianyu Liu, Zhi-	Ruobing Xie, Jie Zhou, Mark Gerstein, Dahai Li,	771
716	wei Jiang, Baobao Chang, Yujia Qin, Wangchunshu	Zhiyuan Liu, and Maosong Sun. 2023b. Toolllm:	772
717	Zhou, Yilun Zhao, Arman Cohan, and Mark Gerstein.	Facilitating large language models to master 16000+	773
718	2023. Ml-bench: Large language models leverage	real-world apis .	774
719	open-source libraries for machine learning tasks .		

775	Samyam Rajbhandari, Olatunji Ruwase, Jeff Rasley,	Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2019.	833
776	Shaden Smith, and Yuxiong He. 2021. Zero-infinity:	Representation learning with contrastive predictive	834
777	Breaking the gpu memory wall for extreme scale	coding.	835
778	deep learning. In <i>Proceedings of the International</i>		
779	<i>Conference for High Performance Computing, Net-</i>	Ke Wang, Houxing Ren, Aojun Zhou, Zimu Lu, Sichun	836
780	<i>working, Storage and Analysis</i> , SC '21, New York,	Luo, Weikang Shi, Renrui Zhang, Linqi Song,	837
781	NY, USA. Association for Computing Machinery.	Mingjie Zhan, and Hongsheng Li. 2023a. Mathcoder:	838
		Seamless code integration in llms for enhanced math-	839
		ematical reasoning.	840
782	Baptiste Rozière, Jonas Gehring, Fabian Gloeckle,	Xiaoxuan Wang, Ziniu Hu, Pan Lu, Yanqiao Zhu, Jieyu	841
783	Sten Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi	Zhang, Satyen Subramaniam, Arjun R. Loomba,	842
784	Adi, Jingyu Liu, Tal Remez, Jérémy Rapin, Artyom	Shichang Zhang, Yizhou Sun, and Wei Wang.	843
785	Kozhevnikov, Ivan Evtimov, Joanna Bitton, Manish	2023b. Scibench: Evaluating college-level scientific	844
786	Bhatt, Cristian Canton Ferrer, Aaron Grattafiori, Wen-	problem-solving abilities of large language models.	845
787	han Xiong, Alexandre Défossez, Jade Copet, Faisal		
788	Azhar, Hugo Touvron, Louis Martin, Nicolas Usunier,	Xingyao Wang, Zihan Wang, Jiateng Liu, Yangyi Chen,	846
789	Thomas Scialom, and Gabriel Synnaeve. 2023. Code	Lifan Yuan, Hao Peng, and Heng Ji. 2023c. Mint:	847
790	llama: Open foundation models for code.	Evaluating llms in multi-turn interaction with tools	848
		and language feedback.	849
791	Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu,	Chenfei Wu, Shengming Yin, Weizhen Qi, Xiaodong	850
792	Junxiao Song, Mingchuan Zhang, Y. K. Li, Y. Wu,	Wang, Zecheng Tang, and Nan Duan. 2023. Visual	851
793	and Daya Guo. 2024. Deepseekmath: Pushing the	chatgpt: Talking, drawing and editing with visual	852
794	limits of mathematical reasoning in open language	foundation models.	853
795	models.		
796	Yongliang Shen, Kaitao Song, Xu Tan, Dongsheng Li,	Qiantong Xu, Fenglu Hong, Bo Li, Changran Hu,	854
797	Weiming Lu, and Yueting Zhuang. 2023. Hugging-	Zhengyu Chen, and Jian Zhang. 2023a. On the	855
798	gpt: Solving ai tasks with chatgpt and its friends	tool manipulation capability of open-source large	856
799	in huggingface. In <i>Advances in Neural Information</i>	language models.	857
800	<i>Processing Systems.</i>		
801	Yifan Song, Weimin Xiong, Dawei Zhu, Wenhao Wu,	Yiheng Xu, Hongjin Su, Chen Xing, Boyu Mi, Qian	858
802	Han Qian, Mingbo Song, Hailiang Huang, Cheng	Liu, Weijia Shi, Binyuan Hui, Fan Zhou, Yitao Liu,	859
803	Li, Ke Wang, Rong Yao, Ye Tian, and Sujian Li.	Tianbao Xie, Zhoujun Cheng, Siheng Zhao, Ling-	860
804	2023. Restgpt: Connecting large language models	peng Kong, Bailin Wang, Caiming Xiong, and Tao	861
805	with real-world restful apis.	Yu. 2023b. Lemur: Harmonizing natural language	862
		and code for language agents.	863
806	Liangtai Sun, Yang Han, Zihan Zhao, Da Ma, Zhennan	Rui Yang, Lin Song, Yanwei Li, Sijie Zhao, Yixiao Ge,	864
807	Shen, Baocai Chen, Lu Chen, and Kai Yu. 2023. Sci-	Xiu Li, and Ying Shan. 2023. Gpt4tools: Teaching	865
808	eval: A multi-level large language model evaluation	large language model to use tools via self-instruction.	866
809	benchmark for scientific research.		
810	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	Da Yin, Faeze Brahman, Abhilasha Ravichander, Khy-	867
811	bert, Amjad Almahairi, Yasmine Babaei, Nikolay	athi Chandu, Kai-Wei Chang, Yejin Choi, and	868
812	Bashlykov, Soumya Batra, Prajwal Bhargava, Shruti	Bill Yuchen Lin. 2023. Lumos: Learning agents	869
813	Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton	with unified data, modular design, and open-source	870
814	Ferrer, Moya Chen, Guillem Cucurull, David Esiobu,	llms.	871
815	Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller,		
816	Cynthia Gao, Vedanuj Goswami, Naman Goyal, An-	Longhui Yu, Weisen Jiang, Han Shi, Jincheng Yu,	872
817	thony Hartshorn, Saghar Hosseini, Rui Hou, Hakan	Zhengying Liu, Yu Zhang, James T Kwok, Zhen-	873
818	Inan, Marcin Kardas, Viktor Kerkez, Madian Khabsa,	guo Li, Adrian Weller, and Weiyang Liu. 2023.	874
819	Isabel Kloumann, Artem Korenev, Punit Singh Koura,	Metamath: Bootstrap your own mathematical ques-	875
820	Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Di-	tions for large language models. <i>ArXiv preprint,</i>	876
821	ana Liskovich, Yinghai Lu, Yuning Mao, Xavier Mar-	abs/2309.12284.	877
822	tinet, Todor Mihaylov, Pushkar Mishra, Igor Moly-		
823	bog, Yixin Nie, Andrew Poulton, Jeremy Reizen-	Lifan Yuan, Yangyi Chen, Xingyao Wang, Yi R. Fung,	878
824	stein, Rashi Rungta, Kalyan Saladi, Alan Schelten,	Hao Peng, and Heng Ji. 2023a. Craft: Customiz-	879
825	Ruan Silva, Eric Michael Smith, Ranjan Subrama-	ing llms by creating and retrieving from specialized	880
826	nian, Xiaoqing Ellen Tan, Binh Tang, Ross Tay-	toolsets.	881
827	lor, Adina Williams, Jian Xiang Kuan, Puxin Xu,		
828	Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan,	Zheng Yuan, Hongyi Yuan, Chengpeng Li, Guanting	882
829	Melanie Kambadur, Sharan Narang, Aurelien Ro-	Dong, Keming Lu, Chuanqi Tan, Chang Zhou, and	883
830	driguez, Robert Stojnic, Sergey Edunov, and Thomas	Jingren Zhou. 2023b. Scaling relationship on learn-	884
831	Scialom. 2023. Llama 2: Open foundation and fine-	ing mathematical reasoning with large language mod-	885
832	tuned chat models.	els.	886

- Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruoqi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, Cong Wei, Botao Yu, Ruibin Yuan, Renliang Sun, Ming Yin, Boyuan Zheng, Zhenzhu Yang, Yibo Liu, Wenhao Huang, Huan Sun, Yu Su, and Wenhua Chen. 2023a. [Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi](#).
- Xiang Yue, Xingwei Qu, Ge Zhang, Yao Fu, Wenhao Huang, Huan Sun, Yu Su, and Wenhua Chen. 2023b. [Mammoth: Building math generalist models through hybrid instruction tuning](#).
- Dan Zhang, Ziniu Hu, Sining Zhoubian, Zhengxiao Du, Kaiyu Yang, Zihan Wang, Yisong Yue, Yuxiao Dong, and Jie Tang. 2024. [Sciglm: Training scientific language models with self-reflective instruction annotation and tuning](#).
- Wenxuan Zhang, Sharifah Mahani Aljunied, Chang Gao, Yew Ken Chia, and Lidong Bing. 2023a. [M3exam: A multilingual, multimodal, multilevel benchmark for examining large language models](#).
- Yifan Zhang, Jingqin Yang, Yang Yuan, and Andrew Chi-Chih Yao. 2023b. [Cumulative reasoning with large language models](#).
- Yilun Zhao, Hongjun Liu, Yitao Long, Rui Zhang, Chen Zhao, and Arman Cohan. 2023. [Knowledgemath: Knowledge-intensive math word problem solving in finance domains](#).
- Aojun Zhou, Ke Wang, Zimu Lu, Weikang Shi, Sichun Luo, Zipeng Qin, Shaoqing Lu, Anya Jia, Linqi Song, Mingjie Zhan, and Hongsheng Li. 2023. [Solving challenging math word problems using gpt-4 code interpreter with code-based self-verification](#).

A Detailed Related Work

A.1 Scientific Reasoning

Scientific reasoning can be roughly categorized into two branches: (1) mathematical reasoning and (2) reasoning across other scientific domains.

Mathematical Reasoning. Mathematical (math) reasoning has attracted much more attentions recently. Thanks to abundant training datasets and corpus, there are intensive studies for more powerful math-oriented LLMs by prompt engineering (Qian et al., 2023; Zhang et al., 2023b; Zhou et al., 2023), instruction-tuning (Yuan et al., 2023b; Yue et al., 2023b; Gou et al., 2023b; Yu et al., 2023; Wang et al., 2023a) and even pre-training (Luo et al., 2023; Azerbayev et al., 2023; Chern et al., 2023). Regarding instruction-tuning, we notice that recent studies have automatically constructed high-quality instructions from GPT-4, *i.e.*, fine-tuning open-source LLMs by Program-of-thought (PoT; Chen et al. 2023a) prompting. It enables open-source LLMs to present remarkable performance, even comparable with GPT-4.

Reasoning across Other Domains. There have been intensive works on scientific LLMs (Bran et al., 2023; Jin et al., 2023; Fang et al., 2023) and benchmarks (Hendrycks et al., 2021a; Huang et al., 2023; Zhang et al., 2023a; Yue et al., 2023a; Sun et al., 2023). However, they primarily target on problems involving less complicated reasoning like knowledge retrieval or simple tool utilization.

Regarding complicated scientific reasoning problems (Chen et al., 2023b; Wang et al., 2023b), questions are scattered among diverse topics and each topic additionally requires domain-specific knowledge. So annotating questions and their solutions domain by domain is much more labor-consuming. Most current benchmarks (Chen et al., 2023b; Wang et al., 2023b; Zhao et al., 2023) merely include hundreds of questions (in all; less for each single domain) from textbooks and provide no training samples. A concurrent work (Zhang et al., 2024) develop a large-scale scientific training corpus, but only focuses three common domains: math, physical and chemistry. Accordingly, the progress of reasoning tasks in these domains is slower than that in math domain: the most competitive approach only achieves 50% and 35% on TheoremQA and SciBench, respectively, not to mention methods built on open-source LLMs. Instead of developing an omniscient and proficient

LLMs on reasoning tasks across various scientific domains, we believe it is more practical to teach LLMs the ability to use domain-specific tools to facilitate their reasoning abilities in some domain when external functions (toolset) are attached.

A.2 Tool Learning

LLMs, both proprietary ones and open-source ones, demonstrate promising capabilities leveraging external tools to solve problems beyond their limits (Qin et al., 2023a). Combined with specific tools, these *tool-augmented* LLMs achieve great success on various tasks such as machine learning (Wu et al., 2023; Shen et al., 2023; Patil et al., 2023; Yang et al., 2023; Liu et al., 2023), question answering (Peng et al., 2023; Gou et al., 2023a), daily assistance (Xu et al., 2023a; Qin et al., 2023b; Song et al., 2023; Gao et al., 2023), *etc.*

Previous work usually pre-defines several tools, *e.g.*, equation solver or calculator, to facilitate math reasoning tasks (Gou et al., 2023a; Lu et al., 2023; Hao et al., 2023; Chen et al., 2023c; Wang et al., 2023c; Xu et al., 2023b; Yin et al., 2023). Cai et al. (2023) generalize the concept of tools to *Program functions*. Following this concept, CREATOR (Qian et al., 2023) scale up the function number towards thousand level. However, these ad-hoc, argument-free functions are more like solution wrapper rather than well-generalized tools. CRAFT (Yuan et al., 2023a) targetedly design an automatic pipeline to extract generalized functions for tool-use. Though leading to improvement, these functions are still not generalized enough and serve more as reference rather than as tools for direct calling. Ouyang et al. 2023 ask LLM to generate chemistry formulae as knowledge reference to assist the following reasoning and achieve enhanced performance on chemistry questions in SciBench. Similar as our attached toolset, Zhao et al. (2023) maintain a *knowledge bank* in which saves more than 900 financial definitions/equations/models as the format of functions for retrieval and use. To our best knowledge, our work is the first which (1) fine-tunes open-source, tool-augmented LLM agents for scientific reasoning tasks and (2) provides a benchmark covering multiple scientific domains to evaluate LLMs’ tool-use abilities.

B Training Details

B.1 Retriever

To fine-tune a retriever, we construct the training samples from MATHFUNC. We concatenate the question and its planning as the query, and view the generated functions as the keys. We finally collect a total of 8603 query-key pairs for training, and split 10% training samples as validation set.

$$\begin{aligned}\text{query} &= [q; G_q] \\ \text{key} &= f \in \tilde{F}_q\end{aligned}$$

We follow DPR (Karpukhin et al., 2020) to train a dense retriever R . We use ROBERTA-BASE (Liu et al., 2019) as the backbone. We set the training step as 500, the batch size as 128 and the learning rate as $2e-5$. We also set the temperature coefficient of the InfoNCE loss (van den Oord et al., 2019) as 0.07. We run this experiment on a single NVIDIA Quadro RTX8000 GPU. The whole training process lasts for about 20 minutes.

B.2 Planning and Action

We fine-tune CodeLlama (Rozière et al., 2023), Mistral (Jiang et al., 2023) and DeepMath (Shao et al., 2024) on MATHFUNC to develop the planning and action modules in our tool-augmented agents SCAGENT series, yielding SCAGENT-CODER, SCAGENT-MISTRAL and SCAGENT-DEEPMATH respectively. We set the global batch size as 128. We use the learning rate as $2e-5$ for CodeLlama, $2e-6$ for Mistral, and $5e-6$ for DeepMath. We use a cosine scheduler with a 3% warm-up period for 2 epochs. We train all models with ZeRO Stage3 (Rajbhandari et al., 2021) on 8 V100 GPUs. The whole training process lasts for about 3 hours for 7B LLMs and 7.5 hours for 13B LLMs.

The planning and action modules share the same model but act differently with different input instructions. We detail the format of planning and action instructions as below:

Planning. Given a question q , we construct a planning sample as $(I_{\text{plan}}(q), G_q)$, where $I_{\text{plan}}(q)$ is the input instruction, G_q is the output, and $I_{\text{plan}}(\cdot)$ is the template for planning module. We provide an example of planning instruction as below:

Listing 1: An example of the planning sample. We separate the input instruction and output answer by the dashed line.

Read the following question and provide a high-level, step-by-step plan for this problem.

Question: Two complementary angles are in a ratio of $\$3:2\$$. What is the measure, in degrees, of the smaller angle?

Answer:

To solve this problem, we need to perform the following steps:
- Set up an equation using the given ratio and the fact that the angles are complementary.
- Solve the equation for the variable representing the ratio.
- Calculate the measure of the smaller angle using the found variable value.

Action. Given a question q , we construct an action sample as $(I_{\text{action}}(q, F_q), S_q)$, where $I_{\text{action}}(q, F_q)$ is the input instruction, S_q is the output answer, and $I_{\text{action}}(\cdot)$ is the template for action module. F_q are retrieved functions from the toolset in MATHFUNC. We adopt the retriever in Appendix B.1 to retrieve three functions in the toolset. We provide an example of action instruction as below:

Listing 2: An example of the action sample. We separate the input instruction and output answer by the dashed line. We only show one retrieved function in this sample for visualization convenience.

Read the following mathematical question and answer it. Please note that you could (optionally, not required) call the following functions when writing the code:

```
def angle_ratio_to_degrees(ratio):
    """
    Computes the angle measures in
    degrees given their ratio.

    Parameters:
    - ratio (tuple): The ratio of the
      angles.

    Returns:
    - tuple: The angle measures in
      degrees.

    Formula:
    - angle_1 = ratio[0] * x
    - angle_2 = ratio[1] * x
    - angle_3 = ratio[2] * x
    - where x is the common variable such
      that angle_1 + angle_2 + angle_3
      = 180 degrees
    """
    sum_ratios = sum(ratio)
    x = 180 / sum_ratios
    return tuple(angle * x for angle in
                  ratio)
```

[...Two more retrieved functions...]

Question: Two complementary angles are in a ratio of 3:2. What is the measure, in degrees, of the smaller angle?

Answer:

```
-----
# Define the ratio of the angles.
ratio = (3, 2)
# Since the angles are complementary,
# they add up to 90 degrees.
# We can find the value of each part of
# the ratio by dividing 90 by the sum
# of the parts of the ratio.
value_of_each_part = 90 / sum(ratio)
# The measure of the smaller angle is
# the value of each part times the
# smaller part of the ratio.
smaller_angle = value_of_each_part * min
(ratio)
# Print the result.
print(smaller_angle)
```

C Evaluation Details

We evaluate (1) the tool-use abilities of LLMs in scientific reasoning task, and (2) their performance when supplemented without any toolsets for comparison. We detail these two settings as below:

With Toolsets. We use the unified PoT-based prompt (Chen et al., 2023a) for all pretraining-based models and our SCIAgent series. The unified prompt consists of a short task description and two demonstrations. We show the prompt in Appendix F.4. For each question, we provide three retrieved functions and instruct LLMs to use them if (and only if) necessary. Note that we use the same retriever, *i.e.*, fine-tuned from MATHFUNC, for all LLMs. For MAMmoTH-Coder and ToRA-Coder which are fine-tuned on specific (tool-agnostic) instructions, we try to enable them to use retrieved tools while keeping the formats of their original instructions as much as possible. Specifically, we append a short *tool-augmented* description at the end of their original prompts:

[original prompt]

Please note that you could (optionally, not required) call the following functions when writing the program:

[retrieved functions]

Without Toolsets. Similar as above, we use the unified PoT-based prompt (Chen et al., 2023a) shown in Appendix F.5 for all pretraining-based models and our SCIAgent series. And we follow the orig-

inal instructions used for MAMmoTH-Coder and ToRA-Coder to evaluate their performance.

D Details of SciTOOLBENCH Annotation

We provide a more thorough description about SciTOOLBENCH construction in this section. This semi-automatic annotation pipeline involves both GPT-4 and humans to balance the quality and cost. Specifically, we enlist two authors to serve as human annotators. Both of them are graduate students with proficiency in English. Additionally, they hold Bachelor of Science and/or Engineering degrees and have completed undergraduate-level courses in the five scientific domains corresponding to our benchmark. We detail the three subsequent sub-modules in our annotation pipeline, *i.e.*, question filtering, positive function construction and negative function construction, as below.

D.1 Question Filtering

We curate the questions from TheoremQA (Chen et al., 2023b) and SciBench (Wang et al., 2023b), both of which are available under the MIT License. Among 1495 questions in these original two datasets, we remove three kinds of questions.

Image-required: There are 37 questions from TheoremQA which include images and necessitate visual understanding abilities. We remove these samples because our benchmark is text-oriented.

Reasoning-agnostic: There are some multi-choice questions from TheoremQA which merely requires the memorization of knowledge points but involves little reasoning process. For example:

Question: The open mapping theorem can be proved by

- (a) Baire category theorem.
- (b) Cauchy integral theorem.
- (c) Random graph theorem.
- (d) None of the above.

We manually check each samples and remove 68 such kind of samples.

Over-difficult: Too hard questions confuse all models and weaken the discrimination of our benchmark. To balance the difficulty and discrimination, we employ 4 advanced proprietary models⁶ to generate related functions and function-augmented program solutions. We generate 6 so-

⁶gpt-4, gpt4-32k, gpt-3.5-turbo, gpt-3.5-turbo-16k

lutions for each model (one generated by greedy decoding and the other five by nucleus sampling with 0.6 temperature) and 24 solutions in all. We view questions that are answered incorrectly by all 24 solutions as *over-difficult* questions. We remove all *over-difficult* questions, and retain 73.5% questions in TheoremQA and 47.8% in SciBench.

By removing three kinds of samples mentioned above, there are a total of 865 questions in our SCITOOLBENCH benchmark.

D.2 Positive Function Construction

Function Generation

In practice, we merge this sub-module to the process of over-difficult question identification. We randomly sample one set of functions which yield correct solutions for each question. As a result, we collect a total of 1216 candidates for the next verification sub-module. We additionally save other functions leading to correct solutions and use them as reference in the refinement sub-module.

Function Verification

We verify the generated functions from both correctness and generalizations. We detail them separately as below.

1. *Correctness*: Since all candidate functions lead to correct solutions, we speculate that almost all of them are correct. We randomly sample 100 functions (20 per domain) and manually check their correctness. The results shown in Table 5 validate our speculation. Therefore, we assume all candidate functions are correct and retain them.

Table 5: The correctness of 100 randomly sampled functions across five domains.

	Correct	Partially Correct	Wrong	All
Math	18	2	0	20
Physics	19	1	0	20
Chemistry	20	0	0	20
Finance	19	0	1	20
EECS	17	3	0	20
All	93	6	1	100

2. *Generalization*: We encounter the similar problem as the function construction in MATHFUNC, i.e., some of the auto-generated functions are not generalized enough. If ad-hoc functions were in the provided toolsets of our benchmark, they would cause a significant overestimation of LLMs’ tool-use abilities. To mitigate it as much as possible, we manually check all candidate functions to ensure their generalization. Specifically, we design

a binary classification task and assign each function a label in {Retained, Refined}. We label a function as *refined* if it had one of the problems listed below: (1) a pure solution wrapper. (2) merely defining a non-generalized expression (likely only occur in this question). (3) the argument names or document describing the special scenario of corresponding question and not being generalized/abstractive enough. (4) including ad-hoc constants or code snippets. The annotators firstly co-annotate 100 functions. We calculate Cohen’s kappa value of their annotation results as 0.85, illustrating an ideal agreement. Therefore, the annotators separately annotate the remaining functions. It takes about 6 hours per annotator to classify about 650 functions. We show some Refined function cases in Figure 10, and the annotation interface in Figure 8.

As a result, we collect 1012 Retained and 206 Refined functions. We keep all Retained as the component of positive functions. We also feed the Refined functions to next refinement sub-module to modify them as much as possible.

Function Refinement

This sub-module aims to rewrite 206 Refined functions to make them qualified. To this end, we associate each function with (1) the question from which it is derived, (2) the function-augmented solutions, and (3) the alternative functions from the generation sub-module (if have). Then we provide them to the annotators. The annotators are asked to rewrite the functions to **improve their generalization** as much as possible. If one function were successfully rewritten, we also require the annotator to write a solution involving the new function to the related question. The solution must yield correct answer to ensure the correctness of the rewritten function. We show some rewritten cases in Figure 10, and the screenshot of the annotation interface in Figure 9.

It takes approximately 12 hours per annotator to check each Refined function and, if applicable, rewrite it. As a consequence, we successfully rewrite 91 Refined functions and drop the remaining ones. We combine these 91 rewritten functions and the 1012 Retained functions to construct 1103 positive functions.

D.3 Negative Function Construction

The positive functions constructed above have satisfied the minimum requirements of the toolset in our benchmark. However, we find that such kind of

benchmark contains shortcuts for LLM to retrieve and use functions. Take a physical question about *frequency-angular conversion* as example, the previous modules construct a positive function named `angular_from_frequency(...)` to solve this question. Without any other similar functions, the LLMs could readily select and use the **only** function by superficial shortcuts. These shortcuts significantly weaken the function-understanding and -use abilities evaluation of our benchmark. To mitigate this problem, we design an additional module to eliminate the shortcuts by constructing some (hard) negative functions for each positive function, like `frequency_from_angular(...)` and `frequency_from_energy(...)` in the above example. Among three similar functions, LLMs are forced to understand their usages and choose proper ones to use. In summary, we add negative functions into the toolset to simulate a more challenging scenario and better evaluate LLMs' tool-use abilities.

Listing 3: Prompt for constructing negative functions

```

Given a function about the {subfield}
field, could you please write two
more functions which satisfy:
- The functions should be in the same
field with the provided function,
while the knowledge point is not
compulsorily the same.
- The functions should be similar, but
not identical with the provided
function.
- The new written functions should be
wrapped as the below format:

New function 1:
```python
[new_written_function_1]
```

New function 2:
```python
[new_written_function_2]
```

```

Specifically, we employ GPT-4 for each positive function to generate two similar but not identical functions as the negative functions. The prompt used is shown as below. We do not validate the correctness of negative functions for simplicity, as they are not intended to be used for any question. We filter the duplicated functions and retain the other 1343 functions in all. By merging the 1103 positive functions and 1343 negative functions, we finally collect a total of 2446 functions in our toolset.

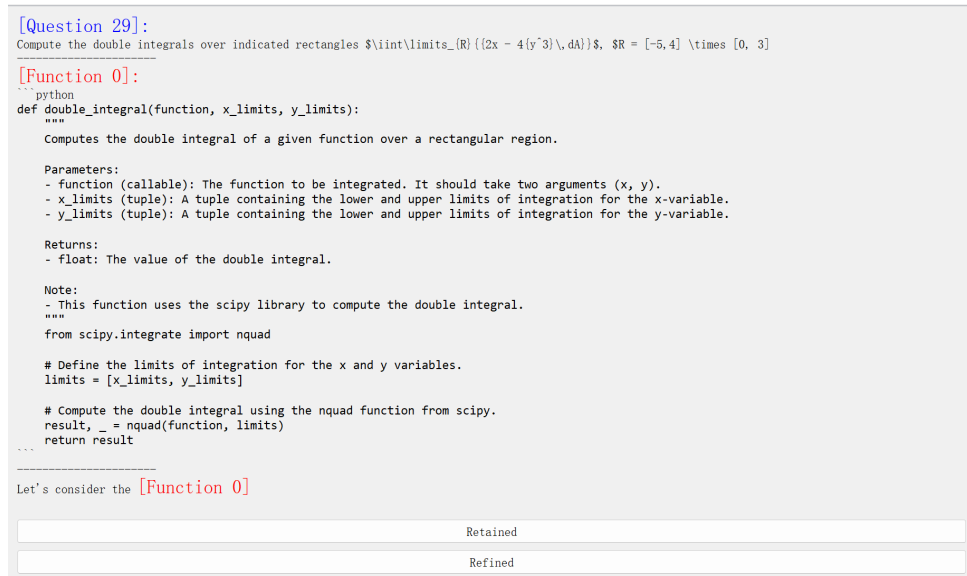


Figure 8: The screenshot of our annotation interface to evaluate functions' generalization.

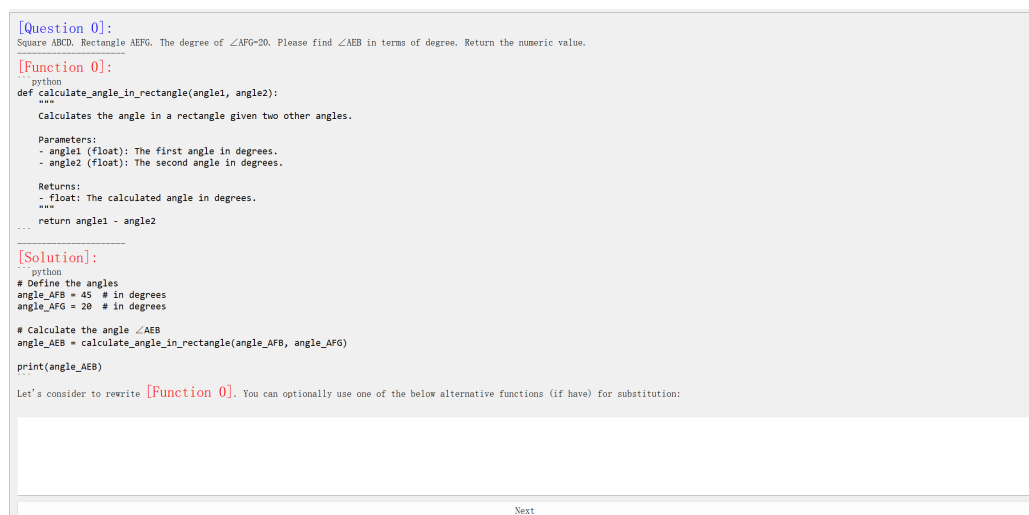


Figure 9: The screenshot of our annotation interface to rewrite functions. We provide no alternative functions in this example for convenience of visualization.

Function before rewriting

```
def birge_vieta(p, tol=1e-3, max_iter=100):
    """
    Finds a real root of the polynomial  $x^3 - 11x^2 + 32x - 22$  using the Birge-Vieta method.

    Parameters:
    - p (float): The initial guess for the root.
    - tol (float, optional): The desired tolerance for the root. Default is 1e-3.
    - max_iter (int, optional): The maximum number of iterations. Default is 100.

    Returns:
    - float: The real root of the polynomial found using the Birge-Vieta method.
    """
    for _ in range(max_iter):
        p_new = p - polynomial(p) / polynomial_derivative(p)
        if abs(p_new - p) < tol:
            return p_new
        p = p_new
    raise ValueError("Birge-Vieta method did not converge within the maximum number of iterations.")
```

Rewrite the specific polynomial (and its derivative) to an argument of the function

Function after rewriting

```
def birge_vieta_iteration(polynomial, p, tol=1e-3, max_iter=100):
    """
    Finds a real root of a polynomial using the Birge-Vieta method.

    Parameters:
    - polynomial (sympy expression): The polynomial for which the root is to be found.
    - p (float): The initial guess for the root.
    - tol (float): The desired tolerance for the root.
    - max_iter (int): The maximum number of iterations allowed.

    Returns:
    - float: The real root of the polynomial, if found within the maximum number of iterations.
    - Raises a ValueError if the root is not found within the maximum number of iterations.
    """
    from sympy import lambdify, diff
    import numpy as np

    # Extract the variable from the polynomial
    variables = list(polynomial.free_symbols)
    if not variables:
        raise ValueError("No variables found in the polynomial.")
    if len(variables) > 1:
        raise ValueError("The polynomial contains more than one variable.")
    variable = variables[0]

    # Compute the derivative of the polynomial
    derivative = diff(polynomial, variable)

    # Convert the polynomial and its derivative to functions
    f = lambdify(variable, polynomial, 'numpy')
    f_prime = lambdify(variable, derivative, 'numpy')

    # Iterate using the Birge-Vieta method
    for _ in range(max_iter):
        p_new = p - f(p) / f_prime(p)
        if np.abs(p_new - p) < tol:
            return p_new
        p = p_new
    raise ValueError("Maximum number of iterations reached without convergence.")
```

Function before rewriting

```
def calculate_emptying_time(height, radius, side_length, g=9.81):
    """
    Calculates the time it takes for a cylindrical tank to go from full to empty.

    Parameters:
    - height (float): The height of the cylindrical tank.
    - radius (float): The radius of the cylindrical tank.
    - side_length (float): The length of the side of the square hole in the bottom of the tank.
    - g (float): The acceleration due to gravity.

    Returns:
    - float: The time it takes for the tank to empty.
    """
    from math import pi, sqrt
    # Calculate the area of the tank and the hole
    tank_area = pi * radius**2
    hole_area = side_length**2

    # Use Torricelli's law to calculate the time
    time = (2 * height * tank_area) / (sqrt(2*g*height) * hole_area)
    return time
```

Function after rewriting

```
def calculate_drain_time(volume, area, gravity=9.81):
    """
    Calculates the time it takes for a cylindrical object to drain using Torricelli's Law.

    Parameters:
    - volume (float): The volume of the cylindrical object.
    - area (float): The area of the hole through which the object is draining.
    - gravity (float): The acceleration due to gravity.

    Returns:
    - float: The time it takes for the object to drain.
    """
    from math import sqrt
    return volume / (area * sqrt(2*gravity))
```

1. Abstract the function description by changing "tank" to "object"
2. Decompose the area calculation and Torricelli's law

Function before rewriting

```
def is_log_concave():
    """
    Determines if the cumulative distribution function (CDF) of the standard Gaussian distribution is log-concave.

    Returns:
    - int: 1 if the CDF is log-concave, 0 otherwise.

    Note:
    - The second derivative of the natural logarithm of the CDF of the standard Gaussian distribution is always non-positive.
    - Therefore, the function is log-concave, and we can return 1 without performing any calculations.
    """
    return 1
```

Rewrite the specific function (and its variable) to an argument of the function

Function after rewriting

```
def is_log_concave(f, x):
    """
    Determines if a given function 'f' with respect to variable 'x' is log-concave.

    Parameters:
    - f (sympy expression): The function for which the log-concavity is to be checked.
    - x (sympy symbol): The variable with respect to which log-concavity is to be checked.

    Returns:
    - bool: True if the function is log-concave, False otherwise.
    """
    from sympy import diff, log, simplify, solveset, S
    from sympy.calculus.util import continuous_domain

    # Compute the first derivative of the logarithm of the function
    first_derivative = diff(log(f), x)

    # Compute the second derivative of the logarithm of the function
    second_derivative = diff(first_derivative, x)

    # Simplify the second derivative
    second_derivative_simplified = simplify(second_derivative)

    # Check if the second derivative is nonpositive (<= 0)
    # We check if the second_derivative is nonpositive for all x in the domain of f
    # For a general solution, we attempt to prove the second derivative is always nonpositive
    return solveset(second_derivative_simplified <= 0, domain=S.Reals) == continuous_domain(second_derivative_simplified, x, S.Reals)
```

Figure 10: Three examples of Refined functions (before rewriting) and their rewritten functions (after rewriting). We also briefly describe the modifications for each example (color in red).

E Examples

Problem & Planing

Read the following mathematical question and answer it. Please note that you could (optionally, not required) call the following functions when writing the code:

Question: Suppose a stock has the following information. It is listed on the London stock exchange and operates throughout Europe. The yield on a UK 10 year treasury is 2.8%. The stock in question will earn 8.6% as per historical data. The Beta for the stock is 1.4, i.e., it is 140% volatile to the changes in the general stock market. What is the expected rate of return?

-
- (1) Calculate the expected return based on the beta and the yield on a UK 10 year treasury.
 - (2) Add the stock's historical return to the expected return.

Retrieved Functions

```
def expected_return(rf, beta, rm):
    """
    Computes the expected return using the Capital Asset Pricing Model (CAPM) formula.

    Parameters:
    - rf (float): The risk-free rate.
    - beta (float): The beta of the portfolio.
    - rm (float): The return on the market.

    Returns:
    - float: The expected return.
    """
    return rf + beta * (rm - rf)

def expected_stock_return(true_prob, u, d):
    """
    Calculates the expected return of the stock in a one-period binomial model.

    Parameters:
    - true_prob (float): The true probability of the stock price going up.
    - u (float): One plus the rate of capital gain on the stock if the price goes up.
    - d (float): One plus the rate of capital loss on the stock if the price goes down.

    Returns:
    - float: The expected return of the stock.
    """
    return true_prob * u + (1 - true_prob) * d

...
```

Solution

```
# Given values.
rf = 0.028 # The yield on a UK 10 year treasury
beta = 1.4 # The stock is 140% volatile to the changes in the general stock market
rm = 0.086 # The stock in question will earn 8.6% as per historical data
# Calculate the expected return .
expected_return = expected_return(rf, beta, rm)
# Print the result.
print(expected_return)
```

Figure 11: An example of direct function call.

Problem & Planing

Read the following mathematical question and answer it. Please note that you could (optionally, not required) call the following functions when writing the code:

Question: The linear density in a rod 8 m long is $12 / \sqrt{x+1}\text{ kg} / \text{m}$, where x is measured in meters from one end of the rod. Find the average density of the rod.

- (1) Integrate the linear density function with respect to x from 0 to 8.
- (2) Divide the result by the length of the rod.

Retrieved Functions

```
def average_value_of_function(f, a, b):
    """
    Calculates the average value of a function `f` on the interval `[a, b]`.
    Parameters: - f (sympy expression): The function for which the average value is to be found.
    - a (float): The lower limit of the interval.
    - b (float): The upper limit of the interval.
    Returns: float: The average value of the function on the interval.
    """
    from sympy import integrate, symbols
    x = symbols('x')
    integral = integrate(f, (x, a, b))
    average_value = integral / (b - a)
    return average_value

def avg_distance(vectors):
    """
    Finds the average distance from the origin that can be reached by adding and subtracting the given
    vectors.
    Parameters: - vectors (list of numpy arrays): The vectors.
    Returns: float: The average distance.
    """
    from itertools import product
    import numpy as np
    total_distance = 0
    count = 0
    for signs in product([-1, 1], repeat=len(vectors)):
        resultant = sum(v * s for v, s in zip(vectors, signs))
        distance = np.linalg.norm(resultant)
        total_distance += distance
        count += 1
    return total_distance / count

...
```

Solution

```
# The linear density is given by 12 / sqrt(x+1) kg/m.
# The average density is the total mass of the rod divided by its length.
# The total mass of the rod is the integral of the linear density times the length of the rod from 0 to
8.
from sympy import symbols, integrate
x = symbols('x')
length = 8
# Define the linear density function.
linear_density = 12 / (x + 1)**0.5
# Calculate the total mass of the rod.
total_mass = integrate(linear_density, (x, 0, length))
# Calculate the average density by dividing the total mass by the length.
average_density = total_mass / length
print(average_density)
```

Figure 12: An example of implicitly using functions as a reference.

F Prompts

F.1 Tool/Function Collection

Listing 4: Prompt used for planning and function collection

Given the question and the reference solution, do the following things:

- Think about what math knowledge points are required to solve this problem step by step.
- write some python one or more functions to abstract the solution. Please note that the functions should be well-documented as much as possible and not too specific (for example, do not write the values in this problem within the functions. Pass them as the function arguments). We hope your written functions could be re-used in anywhere else.

-Instantiate these functions to solve the problem. The last line of your program should be a 'print' command to print the final answer

Here are some examples you may refer to:

Question: There are integers b, c for which both roots of the polynomial $x^2 - x - 1$ are also roots of the polynomial $x^5 - bx - c$. Determine the product bc .

Answer: Let r be a root of $x^2 - x - 1$. Then, rearranging, we have $r^2 = r + 1$.
Multiplying both sides by r and substituting gives $r^3 = 2r + 1$. Repeating this process twice more, we have $r^4 = 2r^2 + r = 2(r + 1) + r = 3r + 2$ and $r^5 = r(3r + 2) = 3r^2 + 2r = 3(r + 1) + 2r = 5r + 3$. Thus, each root of $x^2 - x - 1$ is also a root of $x^5 - 5x - 3$, which gives $bc = 5 \cdot 3 = \boxed{15}$.

Think: To solve this question, we can follow the steps below: (1) Find the roots of the polynomial $x^2 - x - 1$. (2) Substitute them into the polynomial $x^5 - bx - c$ and obtain two equations. (3) Solve the equations.

Functions:

```
def find_roots_of_polynomial(polynomial, variable):
    """
    Finds the roots of a given polynomial using the sympy library.

    Parameters:
    - polynomial (sympy expression): The polynomial whose roots are to be found.
    - variable (sympy symbol): The variable of the polynomial.

    Returns:
    - list: The roots of the polynomial.
    """
    from sympy import solve
    roots = solve(polynomial, variable)
    return roots

def substitute_roots_into_polynomial(roots, polynomial, variable):
    """
    Substitutes the given roots into the polynomial and returns the resulting
    expressions.

    Parameters:
    - roots (list): The list of roots to be substituted into the polynomial.
    - polynomial (sympy expression): The polynomial into which the roots are to be
      substituted.
    - variable (sympy symbol): The variable of the polynomial.

    Returns:
    - list: The resulting expressions after substituting the roots into the
      polynomial.
    """
    return [polynomial.subs(variable, root) for root in roots]
```

```

1436
1437
1438
1439
1440
1441
1442
1443
1444
1445
1446
1447
1448
1449
1450
1451
1452
1453
1454
1455
1456
1457
1458
1459
1460
1461
1462
1463
1464
1465
1466
1467
1468
1469
1470
1471
1472
1473
1474
1475
1476
1477
1478
1479
1480
1481
1482
1483
1484
1485
1486
1487
1488
1489
1490
1491
1492
1493
1494
1495
1496
1497
1498
1499
1500
1501
1502
1503
1504
1505

```function 3
def solve_equations(equations, variables):
 """
 Solves a system of equations for the specified variables using the sympy library.

 Parameters:
 - equations (list of sympy expressions or a single sympy expression):
 The equations to be solved. If solving a single equation, this can be a single
 expression.
 - variables (list of sympy symbols or a single sympy symbol):
 The variables for which the solution is to be found. If solving for a single
 variable, this can be a single symbol.

 Returns:
 - list of dictionaries: Each dictionary represents a solution, with keys being
 the variables and values being their corresponding values.
 If there's only one solution, the list will contain a single dictionary.
 """

 from sympy import solve
 solution = solve(equations, variables, dict=True)
 return solution
```

Solution:
```python
Import required functions and classes from sympy
from sympy import symbols, Eq

Define the variable and the polynomials
x, b, c = symbols('x b c')
polynomial1 = x**2 - x - 1
polynomial2 = x**5 - b*x - c

Find the roots of the first polynomial
roots = find_roots_of_polynomial(polynomial1, x)

Substitute the roots into the second polynomial
resulting_expressions = substitute_roots_into_polynomial(roots, polynomial2, x)

Set up the equations based on the resulting expressions
equations = [Eq(expr, 0) for expr in resulting_expressions]

Solve the system of equations for b and c
solutions = solve_equations(equations, (b, c))
This linear system has only one solution
solution = solutions[0]

Calculate the product bc
product_bc = solution[b] * solution[c]
print(product_bc)
```

---

Question: Medians  $\overline{DP}$  and  $\overline{EQ}$  of  $\triangle DEF$  are perpendicular. If  $DP = 18$  and  $EQ = 24$ , then what is  $DE$ ?
Answer: Point  $G$  is the centroid of  $\triangle DEF$ , so  $DG:GP = EG:GQ = 2:1$ . Therefore,  $DG = \frac{2}{3}(DP) = 12$  and  $EG = \frac{2}{3}(EQ) = 16$ , so applying the Pythagorean Theorem to  $\triangle EGD$  gives us  $DE = \sqrt{EG^2 + GD^2} = \sqrt{16^2 + 12^2} = 20$ .
Think: Given two perpendicular medians in a triangle, we need to perform the following steps: (1) Identify the relationship between the segments of medians and the centroid. (2) Use the ratios provided to determine the lengths of the individual segments from the centroid to the vertices. (3) Use the Pythagorean theorem to determine the length of the side connecting the two vertices from which the medians originate.

Functions:
```function 1

```

```

1506 def median_segments_length(median_length, ratio):
1507 """
1508 Computes the lengths of the segments of a median split by the centroid.
1509
1510 Parameters:
1511 - median_length (float): Total length of the median.
1512 - ratio (tuple): Ratio in which the centroid splits the median. Default is (2,1)
1513 for standard triangles.
1514
1515 Returns:
1516 - tuple: Lengths of the two segments.
1517
1518 Formula:
1519 - segment_1 = ratio[0]/sum(ratio) * median_length
1520 - segment_2 = ratio[1]/sum(ratio) * median_length
1521 """
1522 segment_1 = ratio[0] / sum(ratio) * median_length
1523 segment_2 = ratio[1] / sum(ratio) * median_length
1524 return segment_1, segment_2
1525 ...
1526
1527 ```function 2
1528 def pythagorean_theorem(a, b):
1529 """
1530 Computes the hypotenuse of a right triangle given two legs.
1531
1532 Parameters:
1533 - a, b (float): Lengths of the two legs.
1534
1535 Returns:
1536 - float: Length of the hypotenuse.
1537
1538 Formula:
1539 - c = sqrt(a^2 + b^2)
1540 """
1541 from sympy import sqrt
1542 return sqrt(a**2 + b**2)
1543 ...
1544
1545 Solution:
1546 ```python
1547 # Given values
1548 DP = 18
1549 EQ = 24
1550
1551 # Point G is the centroid.
1552 ratio = (2,1)
1553 # Determine the lengths of the segments split by the centroid
1554 DG, GP = median_segments_length(DP, ratio)
1555 EG, GQ = median_segments_length(EQ, ratio)
1556
1557 # Use Pythagoras theorem to determine DE
1558 DE = pythagorean_theorem(EG, DG)
1559 print(DE)
1560 ```

```

---

## 1562 F.2 Self-rectification

Listing 5: Prompt used for self-rectification

```

1563 {
1564 "role": "system",
1565 "content": "Suppose you are writing Python programs to solve mathematical
1566 questions. Your original program, however, can not run normally for some
1567 reasons. Given the error message returned from the interpreter, please debug
1568 your program carefully and generate the new program. Your generated
1569 correction MUST FOLLOW the following format:\n```function 1\n[the first
1570 created function]\n```\n...\n```function N\n[the N-th created function]\n```\n
1571

```

```

    ```python\n[solution]\n```.\nDO NOT only quote your output as ```python\n[
solution]\n```"
},
{
    "role": "user",
    "content": "Original Program:\n```function 1\ndef solve_quadratic(a, b, c):\n
    from sympy import symbols, Eq, solve\n x = symbols('x')\n quadratic = Eq(a *
x**2 + b * x + c, 0)\n solutions = solve(quadratic, x)\n return solutions\n
``` \n\n```python\ndef find_d_values():\n a, b, c, d = symbols('a b c d')\n eq1
= Eq(6 * a * d, a + d)\n eq2 = Eq((6 * d - 1) * (1 - a), a + d)\n eq3 = Eq(a,
4 * d - 1)\n solution = solve([eq1, eq2, eq3], (a, b, c, d), dict=True)\n
d_values = [sol[d] for sol in solution]\n return d_values\nd_values =
find_d_values()\nprint(d_values)\n``` \n\nError Message:\nFile <string>:15\n
nFile <string>:9, in find_d_values()\n\nNameError: name 'Eq' is not defined"
},
{
 "role": "assistant",
 "content": "Analysis: The error is occurring because the Eq function from the
sympy module is not recognized in the find_d_values function. This is
happening because the ```from sympy import symbols, Eq, solve``` line is only
inside the ```solve_quadratic``` function, so Eq and other imported functions
and classes are only available within the scope of that function. To fix this,
you should claim import statement within the ```find_d_values``` function as
well.\n\nNew Program:\n```function 1\ndef solve_quadratic(a, b, c):\n from
sympy import symbols, Eq, solve\n x = symbols('x')\n quadratic = Eq(a * x**2
+ b * x + c, 0)\n solutions = solve(quadratic, x)\n return solutions\n``` \n\n
```python\ndef find_d_values():\n from sympy import symbols, Eq, solve\n a, b,
c, d = symbols('a b c d')\n eq1 = Eq(6 * a * d, a + d)\n eq2 = Eq((6 * d -
1) * (1 - a), a + d)\n eq3 = Eq(a, 4 * d - 1)\n solution = solve([eq1, eq2,
eq3], (a, b, c, d), dict=True)\n d_values = [sol[d] for sol in solution]\n
return d_values\nd_values = find_d_values()\nprint(d_values)\n```"
},
{
    "role": "user",
    "content": "finished"
},
{
    "role": "user",
    "content": "Original Program:\n```function 1\ndef matrix_square(matrix):\n from
sympy import Matrix\n if not isinstance(matrix, Matrix):\n matrix = Matrix(
matrix)\n return matrix * matrix\n``` \n\n```function 2\ndef solve_equations(
equations, variables): \n from sympy import solve \n solution = solve(
equations, variables, dict=True) \n return solution \n``` \n\n```function 3\
ndef count_solutions(solutions):\n return len(solutions)\n``` \n\n```python\
nfrom sympy import symbols, Eq \na, d = symbols('a d') \nmatrix = [[a, 3],
[-8, d]] \nsquared_matrix = matrix_square(matrix) \nidentity_matrix = [[1, 0],
[0, 1]] \nequations = [Eq(squared_matrix[i][j], identity_matrix[i][j]) for i
in range(2) for j in range(2)] \nsolution = solve_equations(equations, [a, d
]) \nans = count_solutions(solution)\nprint(ans)\n``` \n\nError Message:\nFile
<string>:17\n\nFile <string>:17, in <listcomp>(.0)\n\nTypeError: 'Add'
object is not subscriptable"
},
{
    "role": "assistant",
    "content": "Analysis: The error message is occurring because when you are
creating the list of equations with the list comprehension, squared_matrix[i
][j] is not working as expected. The squared_matrix is a SymPy Matrix object,
and you should use the .row and .col methods.\n\nNew Program:\n```function 1\
ndef matrix_square(matrix):\n from sympy import Matrix\n if not isinstance(
matrix, Matrix):\n matrix = Matrix(matrix)\n return matrix * matrix\n``` \n\n
```function 2\ndef solve_equations(equations, variables): \n from sympy
import solve \n solution = solve(equations, variables, dict=True) \n return
solution \n``` \n\n```function 3\ndef count_solutions(solutions):\n return len(
solutions)\n``` \n\n```python\nfrom sympy import symbols, Eq \na, d = symbols('
a d') \nmatrix = [[a, 3], [-8, d]] \nsquared_matrix = matrix_square(matrix) \
nidentity_matrix = [[1, 0], [0, 1]] \nequations = [Eq(squared_matrix.row(i)[j
], identity_matrix[i][j]) for i in range(2) for j in range(2)] \nsolution =
solve_equations(equations, [a, d]) \nans = count_solutions(solution)\nprint(
ans)\n```"
}

```

```

1643 {
1644 "role": "user",
1645 "content": "finished"
1646 },
1647 {
1648 "role": "user",
1649 "content": "Original Program:\n```function 1\ndef sum_to_product_cos(alpha, beta)
1650 : \n return (cmath.cos(alpha - beta) + cmath.cos(alpha + beta)) / 2\n```\n\n```
1651 function 2\ndef simplify_expression(expression): \n from sympy import
1652 simplify \n return simplify(expression) \n```\n\n```python\nfrom sympy import
1653 symbols, cos\n\nx = symbols('x')\nexpr = cos(x) + cos(3*x) + cos(7*x) + cos
1654 (9*x)\n\n# Apply sum-to-product formula\nexpr1 = 2 * sum_to_product_cos(5*x,
1655 4*x)\nexpr2 = 2 * sum_to_product_cos(5*x, 2*x)\n\n# Simplify the expression\n
1656 nsimplified_expr = simplify_expression(expr1 + expr2)\n\n# Extract the values
1657 of a, b, c, and d and sum them up\na, b, c, d = 4, 1, 3, 5\nans = a + b + c
1658 + d\nprint(ans)\n```\n\nError Message:\nFile <string>:12\n\nFile <string>:2,
1659 in sum_to_product_cos(alpha, beta)\n\nNameError: name 'cmath' is not defined"
1660 },
1661 {
1662 "role": "assistant",
1663 "content": "Analysis: The error message is occurring because the cmath module is
1664 being used in the sum_to_product_cos function but it has not been imported.
1665 You can fix this by importing the cmath module in this function. New Program
1666 :\n```function 1\ndef sum_to_product_cos(alpha, beta):\n import cmath\n
1667 return (cmath.cos(alpha - beta) + cmath.cos(alpha + beta)) / 2\n```\n\n```
1668 function 2\ndef simplify_expression(expression): \n from sympy import
1669 simplify \n return simplify(expression)\n```\n\n```python\nfrom sympy import
1670 symbols, cos\n\nx = symbols('x')\nexpr = cos(x) + cos(3*x) + cos(7*x) + cos
1671 (9*x)\n\n# Apply sum-to-product formula\nexpr1 = 2 * sum_to_product_cos(5*x,
1672 4*x)\nexpr2 = 2 * sum_to_product_cos(5*x, 2*x)\n\n# Simplify the expression\n
1673 nsimplified_expr = simplify_expression(expr1 + expr2)\n\n# Extract the values
1674 of a, b, c, and d and sum them up\na, b, c, d = 4, 1, 3, 5\nans = a + b + c
1675 + d\nprint(ans)\n```\n"
1676 },
1677 {
1678 "role": "user",
1679 "content": "Error message:\nTraceback (most recent call last):\n File, line 9, in
1680 <module>\n exec(\n File '<string>', line 13, in <module>\n File '<string>',
1681 line 3, in sum_to_product_cos\n File, line 356, in __complex__\n return
1682 complex(float(re), float(im))\n File, line 351, in __float__\n raise
1683 TypeError('Cannot convert expression to float')\nTypeError: Cannot convert
1684 expression to float\n"
1685 },
1686 {
1687 "role": "assistant",
1688 "content": "Analysis: The error occurs because you are trying to use cmath.cos
1689 function (which expects float or complex numbers as arguments) with SymPy
1690 symbolic expressions (5x and 4x, 5x and 2x are symbolic expressions). To fix
1691 this, you should use SymPy's cos function instead of cmath.cos function to
1692 work with symbolic expressions.\n\nNew Program:\ndef sum_to_product_cos(alpha,
1693 beta):\n from sympy import cos\n return (cos(alpha - beta) + cos(alpha +
1694 beta)) / 2\ndef simplify_expression(expression): \n from sympy import
1695 simplify \n return simplify(expression) \nfrom sympy import symbols, cos\n\nx
1696 = symbols('x')\nexpr = cos(x) + cos(3*x) + cos(7*x) + cos(9*x)\n\n# Apply
1697 sum-to-product formula\nexpr1 = 2 * sum_to_product_cos(5*x, 4*x)\nexpr2 = 2 *
1698 sum_to_product_cos(5*x, 2*x)\n\n# Simplify the expression\nnsimplified_expr =
1699 simplify_expression(expr1 + expr2)\n\n# Extract the values of a, b, c, and d
1700 and sum them up\na, b, c, d = 4, 1, 3, 5\nans = a + b + c + d\nprint(ans)"
1701 },
1702 {
1703 "role": "user",
1704 "content": "finished"
1705 },
1706 {
1707 "role": "user",
1708 "content": "Original Program:\n```function 1\ndef tan_double_angle(tan_theta):\n
1709 return 2 * tan_theta / (1 - tan_theta**2)\n```\n\n```function 2\ndef
1710 tan_triple_angle(tan_theta):\n return (3 * tan_theta - tan_theta**3) / (1 - 3
1711 * tan_theta**2)\n```\n\n```function 3\ndef solve_polynomial(polynomial, x):\n
1712 from sympy import solve\n return solve(polynomial, x)\n```\n\n```python\nfrom

```

```

sympy import symbols\ntan_theta = symbols('tan_theta')\ntan_2_theta =
tan_double_angle(tan_theta)\ntan_3_theta = tan_triple_angle(tan_theta)\
nequation = tan_theta + tan_2_theta + tan_3_theta\nsimplified_equation =
simplify_expression(equation)\ntan_theta_solutions = solve_polynomial(
simplified_equation, tan_theta)\n\n# Find the solution in the given range (0
< tan_theta < 1)\ntan_theta_value = [solution for solution in
tan_theta_solutions if 0 < solution < 1][0]\nprint(tan_theta_value)\n```\n\
nError Message:\nFile <string>:13\n\nNameError: name 'simplify_expression' is
not defined"
},
{
"role": "assistant",
"content": "Analysis: The error message is occurring because there is no function
named simplify_expression defined in your script. To fix this, you can
define such a function that uses SymPy's simplify function to simplify
expressions.\n\nNew Program:\n```\nfunction 1\ndef simplify_expression(
expression): \n from sympy import simplify \n return simplify(expression)\n
```\n\n```\nfunction 2\ndef tan_double_angle(tan_theta):\n return 2 * tan_theta
/ (1 - tan_theta**2)\n```\n\n```\nfunction 3\ndef tan_triple_angle(tan_theta):\n
return (3 * tan_theta - tan_theta**3) / (1 - 3 * tan_theta**2)\n```\n\n```\n
function 4\ndef solve_polynomial(polynomial, x):\n from sympy import solve\n
return solve(polynomial, x)\n```\n\n```\npython\nfrom sympy import symbols\
ntan_theta = symbols('tan_theta')\ntan_2_theta = tan_double_angle(tan_theta)\
ntan_3_theta = tan_triple_angle(tan_theta)\nequation = tan_theta +
tan_2_theta + tan_3_theta\nsimplified_equation = simplify_expression(equation
)\ntan_theta_solutions = solve_polynomial(simplified_equation, tan_theta)\n\n
# Find the solution in the given range (0 < tan_theta < 1)\ntan_theta_value =
[solution for solution in tan_theta_solutions if 0 < solution < 1][0]\nprint
(tan_theta_value)\n```"
},
{
"role": "user",
"content": "finished"
}
}

```

F.3 Function-augmented Solutions

Listing 6: Prompt used for the generation of function-augmented solutions (cross-retrieval strategy)

```

You will encounter a mathematical problem and are required to write a piece of
Python code to solve this problem.

Now we have a suite of wrapped functions. Take note:
- The newly provided wrapped functions have NOT been verified. They may be
  irrelevant or potentially flawed.
- It's essential that the solution doesn't overly depend on wrapped functions.
  You're welcome to utilize one or more functions from the new set in your solution
  but only after you've determined:
  (1) Their accuracy.
  (2) Their inclusion significantly streamlines the problem-solving approach.

Additionally take note that
  (1) The last line of your written code shall be a 'print' command to print the
    final answer.
  (2) The wrapped functions should not be duplicated within your code. Instead,
    call them directly if needed.
  (3) Should you need to create custom functions, do so without adding
    documentation comments for the sake of brevity.
  (4) Write simple but clear annotations interleaving your code solution.

"""
Retrieved functions:
[List of called function names from the new set]

```python
[Your Written Python Code.]
```
"""

```

```

1780 For example:
1781 ---
1782 Question: What is the 100th digit to the right of the decimal point in the decimal
1783 representation of  $\frac{13}{90}$ ?
1784
1785 New provided functions:
1786 ```New Function 0
1787 def decimal_representation(numerator, denominator, max_digits=1000):
1788     """
1789     Computes the decimal representation of a fraction.
1790
1791     Parameters:
1792     - numerator (int): The numerator of the fraction.
1793     - denominator (int): The denominator of the fraction.
1794     - max_digits (int): The maximum number of decimal digits to compute.
1795
1796     Returns:
1797     - str: The decimal representation of the fraction as a string.
1798     """
1799
1800     result = ""
1801     remainder = numerator % denominator
1802     for _ in range(max_digits):
1803         remainder *= 10
1804         result += str(remainder // denominator)
1805         remainder %= denominator
1806         if remainder == 0:
1807             break
1808     return result
1809 ...
1810
1811 ```New Function 1
1812 def decimal_to_scientific(decimal_number):
1813     from sympy import log, floor
1814     exponent = -floor(log(decimal_number, 10))
1815     coefficient = decimal_number * 10**(-exponent)
1816     return coefficient, exponent
1817 ...
1818
1819 ```New Function 2
1820 def repeating_decimal_representation(numerator, denominator):
1821     """
1822     Computes the repeating decimal representation of a fraction.
1823
1824     Parameters:
1825     - numerator (int): The numerator of the fraction.
1826     - denominator (int): The denominator of the fraction.
1827
1828     Returns:
1829     - str: The repeating decimal representation of the fraction as a string.
1830     """
1831
1832     # Initialize the result string and a dictionary to store remainders.
1833     result = ""
1834     remainders = {}
1835
1836     # Perform long division to find the decimal representation.
1837     while numerator != 0:
1838         # If the remainder has been seen before, we found the repeating block.
1839         if numerator in remainders:
1840             start = remainders[numerator]
1841             return result[:start] + "(" + result[start:] + ")"
1842         # Otherwise, store the remainder and continue the division.
1843         remainders[numerator] = len(result)
1844         numerator *= 10
1845         result += str(numerator // denominator)
1846         numerator %= denominator
1847
1848     return result
1849 ...

```

```

1850
1851
1852
1853
1854
1855
1856
1857
1858
1859
1860
1861
1862
1863
1864
1865
1866
1867
1868
1869
1870
1871
1872
1873
1874
1875
1876
1877
1878
1879
1880
1881
1882
1883
1884
1885
1886
1887
1888
1889
1890
1891
1892
1893
1894
1895
1896
1897
1898
1899
1900
1901
1902
1903
1904
1905
1906
1907
1908
1909
1910
1911
1912
1913
1914
1915
1916
1917
1918
1919

'''New Function 3
def nth_digit_of_decimal_representation(numerator, denominator, n):
    """
    Computes the nth digit after the decimal point of the decimal representation of a
    fraction.

    Parameters:
    - numerator (int): The numerator of the fraction.
    - denominator (int): The denominator of the fraction.
    - n (int): The position of the digit after the decimal point.

    Returns:
    - int: The nth digit after the decimal point of the decimal representation of the
    fraction.
    """

    # Get the repeating decimal representation of the fraction.
    decimal_representation = repeating_decimal_representation(numerator, denominator)

    # Remove the parentheses from the repeating block.
    decimal_representation = decimal_representation.replace("(", "").replace(")", "")

    # Calculate the nth digit using the repeating block.
    return int(decimal_representation[(n - 1) % len(decimal_representation)])
'''

Retrieved functions:
[decimal_representation, nth_digit_of_decimal_representation]

'''python
# Use the nth_digit_of_decimal_representation function to find the 100th digit
numerator = 13
denominator = 90
n = 100

# Call the function and print the result
result = nth_digit_of_decimal_representation(numerator, denominator, n)
print(result)
'''

---
Question: The square root of  $x^2$  is greater than 3 and less than 4. How many integer
values of  $x^2$  satisfy this condition?

New provided functions:
'''New Function 0
def solve_square_root_equation(a, b, c):
    """
    Solves a square root equation of the form  $\sqrt{ax - b} = c$ .

    Parameters:
    - a (float): Coefficient of x inside the square root.
    - b (float): Constant term inside the square root.
    - c (float): Constant term on the right side of the equation.

    Returns:
    - float: The value of x that satisfies the equation.

    Formula:
    -  $x = (c^2 + b) / a$ 
    """
    return (c**2 + b) / a
'''

'''New Function 1
def find_integer_square_less_than_double():
    """
    Finds the only integer whose square is less than its double.

```

```

1920
1921 Returns:
1922 - int: The integer that satisfies the condition.
1923
1924 Method:
1925 - Iterate through integers starting from 1, and check if the square of the
1926   integer is less than its double.
1927 - If the condition is satisfied, return the integer.
1928 - If the condition is not satisfied for any integer up to a certain limit, return
1929   None.
1930 """
1931 limit = 100
1932 for x in range(1, limit):
1933     if x**2 < 2*x:
1934         return x
1935 return None
1936 ...
1937
1938 ```New Function 2
1939 def solve_equation():
1940     """
1941     Solves the equation  $(x-2)^{(25-x^2)} = 1$  for integer solutions.
1942
1943     Returns:
1944     - list: A list of integer solutions for x.
1945     """
1946     solutions = []
1947
1948     # Case 1: Exponent is 0 ( $25 - x^2 = 0$ )
1949     x1 = 5
1950     x2 = -5
1951     solutions.extend([x1, x2])
1952
1953     # Case 2: Base is 1 ( $x - 2 = 1$ )
1954     x3 = 3
1955     solutions.append(x3)
1956
1957     # Case 3: Base is -1 and exponent is even ( $x - 2 = -1$  and  $25 - x^2 = 2n$  for some
1958       integer n)
1959     x4 = 1
1960     solutions.append(x4)
1961
1962     return solutions
1963 ...
1964
1965 ```New Function 3
1966 def count_integers_in_range(lower_bound, upper_bound, exclude_zero=True):
1967     """
1968     Counts the number of integers within a given range.
1969
1970     Parameters:
1971     - lower_bound (int): The lower bound of the range.
1972     - upper_bound (int): The upper bound of the range.
1973     - exclude_zero (bool): Whether to exclude 0 from the count. Default is True.
1974
1975     Returns:
1976     - int: The number of integers within the range.
1977     """
1978     count = upper_bound - lower_bound + 1
1979     if exclude_zero and lower_bound <= 0 and upper_bound >= 0:
1980         count -= 1
1981     return count
1982 ...
1983
1984 Retrieved functions:
1985 []
1986
1987 ```python
1988 # The lower and upper bounds of x for which  $\sqrt{x} > 3$  and  $\sqrt{x} < 4$ 
1989 lower_bound = 9

```

```

upper_bound = 16

# Counting the number of integers between 9 (exclusive) and 16 (exclusive)
num_integers = len([x for x in range(lower_bound + 1, upper_bound)])

# Printing the result
print(num_integers)

```

F.4 Evaluation with Toolsets

Listing 7: Prompt used for evaluation (setting with toolsets)

```

Read the following questions and answer them. For each question, you are required to
write a Python program to solve it.
Please note that we provide you several functions for each question. You could (
optionally, not required) call the functions to help you to solve the question
if necessary.
Note that the last line of your program should be a 'print' command to print the
final answer

-----
Question:
What is the 100th digit to the right of the decimal point in the decimal
representation of  $\frac{13}{90}$ ?

Functions:
def repeating_decimal_representation(numerator, denominator):
    """
    Computes the repeating decimal representation of a fraction.

    Parameters:
    - numerator (int): The numerator of the fraction.
    - denominator (int): The denominator of the fraction.

    Returns:
    - str: The repeating decimal representation of the fraction as a string.
    """

    # Initialize the result string and a dictionary to store remainders.
    result = ""
    remainders = {}

    # Perform long division to find the decimal representation.
    while numerator != 0:
        # If the remainder has been seen before, we found the repeating block.
        if numerator in remainders:
            start = remainders[numerator]
            return result[:start] + "(" + result[start:] + ")"
        # Otherwise, store the remainder and continue the division.
        remainders[numerator] = len(result)
        numerator *= 10
        result += str(numerator // denominator)
        numerator %= denominator

    return result

def nth_digit_of_decimal_representation(numerator, denominator, n):
    """
    Computes the nth digit after the decimal point of the decimal representation of a
    fraction.

    Parameters:
    - numerator (int): The numerator of the fraction.
    - denominator (int): The denominator of the fraction.
    - n (int): The position of the digit after the decimal point.

    Returns:

```

```

2057     - int: The nth digit after the decimal point of the decimal representation of the
2058         fraction.
2059     """
2060
2061     # Get the repeating decimal representation of the fraction.
2062     decimal_representation = repeating_decimal_representation(numerator, denominator)
2063
2064     # Remove the parentheses from the repeating block.
2065     decimal_representation = decimal_representation.replace("(", "").replace(")", "")
2066
2067     # Calculate the nth digit using the repeating block.
2068     return int(decimal_representation[(n - 1) % len(decimal_representation)])
2069
2070
2071 def decimal_representation(numerator, denominator, max_digits=1000):
2072     """
2073     Computes the decimal representation of a fraction.
2074
2075     Parameters:
2076     - numerator (int): The numerator of the fraction.
2077     - denominator (int): The denominator of the fraction.
2078     - max_digits (int): The maximum number of decimal digits to compute.
2079
2080     Returns:
2081     - str: The decimal representation of the fraction as a string.
2082     """
2083
2084     result = ""
2085     remainder = numerator % denominator
2086     for _ in range(max_digits):
2087         remainder *= 10
2088         result += str(remainder // denominator)
2089         remainder %= denominator
2090         if remainder == 0:
2091             break
2092     return result
2093
2094
2095 Solution:
2096 # find the 100th digit.
2097 numerator = 13
2098 denominator = 90
2099 n = 100
2100
2101 # Call the function and print the result.
2102 result = nth_digit_of_decimal_representation(numerator, denominator, n)
2103 print(result)
2104
2105
2106 -----
2107 Question:
2108 The square root of  $x$  is greater than 3 and less than 4. How many integer values of
2109  $x$  satisfy this condition?
2110
2111 Functions:
2112 def count_integers_in_range(lower_bound, upper_bound, exclude_zero=True):
2113     """
2114     Counts the number of integers within a given range.
2115
2116     Parameters:
2117     - lower_bound (int): The lower bound of the range.
2118     - upper_bound (int): The upper bound of the range.
2119     - exclude_zero (bool): Whether to exclude 0 from the count. Default is True.
2120
2121     Returns:
2122     - int: The number of integers within the range.
2123     """
2124     count = upper_bound - lower_bound + 1
2125     if exclude_zero and lower_bound <= 0 and upper_bound >= 0:
2126         count -= 1

```

```

return count
2127
2128
2129
def find_integer_square_less_than_double():
2130
2131
2132
2133
2134
2135
2136
2137
2138
2139
2140
2141
2142
2143
2144
2145
2146
2147
2148
2149
2150
2151
2152
2153
2154
2155
2156
2157
2158
2159
2160
2161
2162
2163
2164
2165
2166
2167
2168
2169
2170
2171
2172
2173
2174
2175
2176
2177
"""
Finds the only integer whose square is less than its double.

Returns:
- int: The integer that satisfies the condition.

Method:
- Iterate through integers starting from 1, and check if the square of the
  integer is less than its double.
- If the condition is satisfied, return the integer.
- If the condition is not satisfied for any integer up to a certain limit, return
  None.
"""
limit = 100
for x in range(1, limit):
    if x**2 < 2*x:
        return x
return None

def solve_square_root_equation(a, b, c):
"""
Solves a square root equation of the form sqrt(ax - b) = c.

Parameters:
- a (float): Coefficient of x inside the square root.
- b (float): Constant term inside the square root.
- c (float): Constant term on the right side of the equation.

Returns:
- float: The value of x that satisfies the equation.

Formula:
- x = (c^2 + b) / a
"""
return (c**2 + b) / a

Solution:
# We need to find the integer values of x for which sqrt(x) > 3 and sqrt(x) < 4. To
  this end, we can count the number of integers in the range (9, 16) directly.
result = 0
for x in range(10, 16):
    if 9 < x < 16:
        result += 1
print(result)

```

F.5 Evaluation without Toolsets

Listing 8: Prompt used for evaluation (setting without toolsets)

```

Read the following questions and answer them. For each question, you are required to
  write a Python program to solve it.
Please note that we provide you several functions for each question. You could (
  optionally, not required) call the functions to help you to solve the question
  if necessary.
Note that the last line of your program should be a 'print' command to print the
  final answer

-----
Question:
What is the 100th digit to the right of the decimal point in the decimal
  representation of  $\frac{13}{90}$ ?
Solution:

```

```

2194 from decimal import Decimal, getcontext
2195
2196 # Set the precision to 101 (100 digits after decimal + 1 digit before decimal)
2197 getcontext().prec = 101
2198
2199 # Calculate the decimal representation of 13/90
2200 dec = Decimal(13) / Decimal(90)
2201
2202 # Convert the decimal to a string
2203 dec_str = str(dec)
2204
2205 # Get the 100th digit to the right of the decimal point
2206 digit_100th = dec_str[101]
2207
2208 print(digit_100th)
2209
2210 -----
2211 Question:
2212 The square root of  $x$  is greater than 3 and less than 4. How many integer values of
2213  $x$  satisfy this condition?
2214
2215 Solution:
2216 # We need to find the integer values of  $x$  for which  $\sqrt{x} > 3$  and  $\sqrt{x} < 4$ . To
2217 this end, we can count the number of integers in the range (9, 16) directly.
2218 result = 0
2219 for x in range(10, 16):
2220     if 9 < x < 16:
2221         result += 1
2222 print(result)
2223

```
