

SELF-DIRECTED DISCOVERY: HOW LLMs EXPLORE AND OPTIMIZE CONFIGURABLE LLM SOLUTIONS

Anonymous authors

Paper under double-blind review

ABSTRACT

LLM solutions are increasingly replacing specialized machine learning models across various industry domains. While offering simplicity and maintainability advantages, optimizing these workflows remains heavily dependent on expert-driven experimentation. In this paper we test off-the-shelf powerful LLMs for the task of automatically building and iteratively improving LLM-based solutions. We propose a configuration-driven framework which defines such workflows by specifying model parameters, prompt templates, and data transformations. We show that this standardized representation enables automatic iterative improvement loops via optimization agents which are themselves defined within the same framework. When evaluated on challenging datasets, it discovers improved solutions while maintaining interpretability and human verifiability. The improvement loop we instantiate is generic and minimal (using prompts that have not been engineered) and self-referential (improved solutions are discovered with improved documentation and examples). The proposal is a self-improving system that bridges the gap between generic code-generating automatic optimization and more narrowly-focused techniques such as prompt engineering.

1 INTRODUCTION

Applications across domains such as financial services, healthcare, customer support, content moderation, and many others, increasingly rely on LLM solutions where the same inference pattern is applied to thousands or millions of data points. Such unified LLM-based solutions are increasingly preferred over specialized ML solutions, offering the appeal of a single, versatile approach that can handle diverse tasks without requiring domain-specific model training and maintenance (Chkribene et al., 2024; Raza et al., 2025; Saleh et al., 2025). While these LLM inference workflows may initially lack the efficiency of highly optimized, custom machine learning (ML) solutions developed over years, maintainable systems often outweigh the performance gains of complex, specialized alternatives. However, unlike conversational AI systems where each interaction is unique, these LLM inference scenarios involve repeatedly executing identical prompt templates with only the input data varying, creating additional opportunities for optimization, for both quality and other performance goals. While it is becoming commonplace to utilize techniques such as prompt caching (Xiao et al., 2023; Zhu et al., 2023), automatic prompt engineering ((Ramnath et al., 2025; Zhou et al., 2022; Li et al., 2025a; Debnath et al., 2025)) or model selection (Wang et al., 2023; Tanaka et al., 2023; Tang et al., 2024), optimizing repeated inference patterns *holistically* still heavily relies on expert-driven experimentation (Li et al., 2024).

On the other hand, generative models are also transforming ML *research* by acting as powerful assistants that automate routine tasks, augment human creativity, and accelerate the pace of discovery. While LLMs are not yet capable of fully replacing human ML scientists, they significantly enhance efficiency by automating coding, data analysis, hypothesis generation, and manuscript writing (Tang et al., 2025; Wang & et al., 2025; Yang et al., 2024b; Schramowski et al., 2025). In this paper we take a step further in the direction of using generative models to aid in the development of ML solutions by showing that LLMs can successfully drive the optimization of LLM inference workflows and help discover improved solutions. This proposal addresses the gap between generic code-generating automatic optimization and more narrowly-focused techniques such as prompt engineering.

The paper makes the following contributions:

- We address a specific class of ML solutions based on LLM inference and propose a configuration-driven framework that supports *interpretability*. We demonstrate that this constrained yet expressive solution space enables powerful off-the-shelf LLMs to autonomously create and improve workflows across various tasks with minimal task-specific guidance. The generated workflows execute without errors over 95% of the time, achieve task performance that matches standardized prompts on average, and consistently discover significantly better solutions.
- We further implement a simple automatic iterative improvement loop, where other LLMs improve the solutions for custom goals such efficiency or quality. We propose a very simple linear iterative improvement loop using an Analyzer and an Improver agent, and show it can find better solutions to various challenging data sets. Since the Analyzer and Improver are in themselves LLM-based, we define them within the same framework and can thus be the object of optimization as well.
- We prioritize documentation-driven optimization over prompt engineering, which requires extensive manual tuning for each task and optimization objective. Instead of engineered meta-prompts, our LLMs receive minimal generic prompts alongside rich contextual artifacts: framework documentation and workflow examples. This approach enables robust cross-task generalization and sets the stage for a self-evolving system where accumulated knowledge artifacts improve performance across diverse domains without task-specific prompt crafting.

2 CONFIGURATION-DRIVEN LLM AGENT DEFINITION

In this section we test whether off-the-shelf LLMs can generate high-performing agents for new tasks given a well-defined agent specification framework and examples of agents implemented using it. To achieve this we propose a configuration-driven LLM agent framework which promotes interpretability (a human user can easily understand and verify the generated agent) and enables systematic optimization (another LLM can iteratively improve on these agents).

LLM Agent We define an LLM agent as a workflow that makes a *single* LLM inference call, optionally preceded by input pre-processing and followed by output processing.¹ More precisely, let $\mathcal{D} = \{(x_1, y_1), \dots, (x_n, y_n)\}$ denote a dataset where each data point consists of input x_i and optional ground truth y_i . An LLM agent A is defined by:

- A prompt template T containing fixed text and placeholder variables
- A set of input functions $\{f_1, f_2, \dots, f_k\}$ that transform input data x_i
- An LLM model (θ_{LLM}) alongside inference hyper-parameters ρ (temperature, max tokens, etc.)

For input x_i , the agent generates a prompt p_i and then output o_i as :

$$p_i = T(f_1(x_i, c), \dots, f_k(x_i, c), c) \quad \text{and} \quad o_i = A(x_i) = \theta_{\text{LLM}}(p_i; \rho)$$

where c stands for additional context (examples, external data sources, etc.). Dataset-level performance is computed as $\bar{E} = \frac{1}{n} \sum_{i=1}^n E(o_i, x_i, y_i)$ where E is a task-appropriate evaluation function.

Agent Implementation At the level of implementation, we describe an LLM agent via three files: Configuration (json), Agent Class Implementation (Python) and Input Schema File (json). See Appendix B for more details. The *configuration* file specifies the agent’s behavior: model settings, prompt template, input processing functions, and output format. The agent class implements the pre-processing functions referenced in the configuration. These functions handle data transformations, feature extraction, example retrieval, and other per-datapoint operations. The input schema file defines the agent’s interface with other components, specifying what data the agent can access and how it communicates with upstream and downstream tasks. This three-file structure enforces clear separation between declarative configuration, implementation logic, and interface contracts, enabling systematic optimization while maintaining interpretability.

At run-time, agents process data points sequentially through a standardized pipeline: load and validate configuration files, apply input transformations, render the prompt template, call the LLM,

¹This definition differs from the broader usage of “agent” in the literature, which typically encompasses multi-step reasoning, tool usage, memory systems, and iterative planning. Our constrained definition focuses on the fundamental building block of LLM-based solutions.

and parse the response according to the specified output format. Each agent is constrained to exactly one LLM inference call—multi-step reasoning requiring multiple LLM interactions must be orchestrated at the workflow level using separate agents.

Modifiability Control To guide other LLMs in improving task agents, we introduce modifiability control. Each configuration section includes a flag *modifiable* $\in$ {true, false} that helps meta-agents understand which sections can be modified. For example, `model.modifiable = true` allows modification of model parameters, while `output.modifiable = false` preserves the response schema. This enables systematic exploration while maintaining configurable constraints based on the specific improvement loop requirements—for example, preserving output schemas for downstream compatibility allowing flexible optimization across all components.

Appendix C shows a complete math problem solving agent example and Appendix D shows how this design supports diverse implementations including ML integration, in-context learning, and dynamic input transformation.

2.1 LLM-BASED AGENT GENERATION

To evaluate the configuration-driven framework, we first examine whether LLMs can generate valid task agents A_t given only framework documentation and minimal task context. This evaluation aims to validate the framework’s clarity and usability, as successful agent generation from documentation and minimal context indicates that the abstractions are sufficiently well-designed for automated solution synthesis.

2.1.1 EXPERIMENTAL SETUP

We evaluate agent generation capabilities across the following problems:

AIME2024 and AIME 2025 contain 30 competition-level mathematics problems each, requiring integer answers between 000-999 (MAA Committees). These challenging problems span algebra, geometry, number theory, and combinatorics, requiring complex mathematical reasoning that remains difficult for current (non-reasoning) models. We evaluate exact match accuracy and use AIME2024 Part I for development, with remaining datasets for testing.

Math500: 500 competition mathematics problems with detailed solutions across 7 subjects (Algebra, Counting & Probability, Geometry, Intermediate Algebra, Number Theory, Precalculus, Prealgebra) (Hendrycks et al., 2021; Lightman et al., 2024). For Math500, evaluation uses mathematical expression normalization to account for equivalent mathematical representations². We sample 100 data points for development and leave the remaining 400 for test.

MMLU-Pro is a benchmark designed to challenge LLMs beyond the original MMLU (Wang et al., 2024). It features reasoning-focused multiple-choice questions across 14 domains with ten answer options instead of four, creating a more demanding evaluation. We use validation data for development and sample 200 test problems.

We test if LLMs can generate working agents for these tasks using agent framework documentation and minimal task information. Specifically we implement the agent generation using a single meta-agent defined in the same framework. Table 1 shows the prompt template used by this meta-agent, where the placeholders stand for:

1. **Framework Documentation** and **Agent Example**: The complete LLM Agent README (in supplementary material, totaling 5k tokens) and a single complete agent example. For all tasks we use the basic AIME2024 agent shown in Appendix C. Note that this example is out-of-domain for MMLU-Pro.
2. **Task Data Point** and **Task description**: Task data points are randomly drawn from a different data split, with the purpose of exemplifying the task to be solved and the input/output format. Task descriptions are 1-2 sentences describing the task objective (see Table 1).

²We use the grader at <https://github.com/openai/prm800k/blob/main/prm800k/grading/grader.py>

Table 1: Prompt template used by the agent generating (meta-)agent, which generates task agents A_t using the task information in this table, alongside the [Framework Documentation](#) and [Agent Example](#) described above.

Dataset	Task Data Point	Task Description	Generator Prompt Template
AIME2024	Every morning Aya goes for a 9-kilometer-long walk and stops at a coffee shop afterwards. When [...] Find the number of minutes the walk takes her, including the t minutes spent in the coffee shop. Answer: 204 Solution: <i>detailed solution</i>	AIME (American Invitational Mathematics Examination) problems require integer answers between 000 and 999. The agent must show detailed reasoning and provide the final numerical answer.	Instructions: You are an expert AI agent developer. Task: Task Description Requirements: Generate exactly 3 files: (1) Agent configuration JSON, (2) Agent implementation Python, (3) Input schema JSON Schema. The agent must accept problem data as input, generate detailed reasoning, output correct format, handle parsing errors gracefully, and follow framework patterns. Framework Doc: Framework Documentation Agent Example: Agent Example Sample Data: Task Data Point Output Format: JSON only with agent.config, agent.code, input.schema fields. Generate the agent now:
AIME2025	Find the sum of all integer bases $b > 9$ for which 17_b is a divisor of 97_b . Answer: 070		
Math500	Simplify $\tan 100^\circ + 4 \sin 100^\circ$. Answer: $-\sqrt{3}$ Subject: Precalculus, Level: 2	MATH problems require answers in LaTeX format within <code>\boxed</code> tags. The agent must show detailed reasoning and provide the final answer in the correct format.	
MMLU-Pro	Which of the following represents an accurate statement concerning arthropods? A. They possess an exoskeleton composed primarily of peptidoglycan. B. They possess an open circulatory [...] C. [...] Answer: B	MMLU-Pro problems are multiple-choice questions across academic subjects (biology, math, physics, etc.) with 10 answer choices (A-J). The agent must provide detailed step-by-step reasoning and select the correct answer choice.	

The LLM must infer an appropriate agent structure, input processing logic, prompt design, and output schema solely from the framework documentation, complete agent example and data sample; no additional examples of agent configurations, code samples, or task-specific guidance are provided. However in the future we see this as a self-evolving framework, where the amount of artifacts increases with every problem solved, creating better and richer context for the LLM or code assistant.

Evaluation and baselines In order to facilitate comparisons, we standardize Claude 3.5 Sonnet v2 for all task agents by removing other model references from the framework documentation. The meta-agent uses the same model with temperature 0.7 for exploration.

We run the meta-agent 10 times per task and perform static validation on generated configurations (structure, model names, required fields). Runtime issues include undefined placeholders (non-fatal) and syntax/execution errors (fatal). Success rate measures the percentage of agents executing without fatal errors on development data. For successful agents, we further evaluate task performance using multiple runs due to output variability: 5 runs for AIME datasets and 2 runs for others. We report accuracy averaged across all runs (Pass@1).

We compare against standard prompts for each task: basic CoT for AIME datasets, CoT with one example for Math500, and 5-shot CoT for MMLU-Pro. Exact prompts and evaluation methodology follow <https://www.vals.ai/benchmarks/>.

While the framework supports building sophisticated agents beyond prompt variations, we expect limited use of its full expressive power in these experiments given the minimal framework documentation and single agent example provided.

2.1.2 RESULTS

Results are shown in Table 2.

The meta-agent successfully generated executable agents across all datasets with 100% success rate, even for MMLU-Pro—a previously unseen task provided with only a single example and brief description. The generated agents demonstrated strong performance with median performance approaching that of the standard prompts for the tasks. On all tasks, the best agents significantly out-

Table 2: Performance when generating 10 agents for each task: Success rates (percentage of generated agents that execute without errors) and task performance metrics (median and max). Base accuracy reports performance of standard prompts used for these tasks (see <https://www.vals.ai/>)

Dataset	Success Rate	Base Acc.	1 Task data point		5 Task data points	
			Median Acc.	Max Acc.	Median Acc.	Max Acc.
AIME24 I	100%	22.7	21.3	25.3	20.7	25.3
Math500	100%	70.0	66.0	74.0	69.0	73.0
MMLU-Pro	100%	82.0	82.1	87.1	82.9	90.0

performed the baselines, reaching maximum accuracies of 25% on AIME2024, 74% on Math500, and 90% on MMLU-Pro (an 8% absolute improvement). The primary challenge for agent generation was ensuring consistency between the output format specification in the configuration, the formatting instructions in the prompt template, and the evaluation requirements (which expect the same output fields as shown in the task examples). Misalignment across these components resulted in task agents that produced un-parseable outputs, leading to evaluation failures and zero accuracy scores.

All agents implemented dataset-specific preprocessing and incorporated detailed chain-of-thought reasoning through structured step-by-step problem breakdown. The AIME2024 best agent wrote a 6-step processes: “1. Read and understand the problem carefully, 2. Break down key information, 3. Plan your solution [...] 6. Verify answer is integer between 000-999”. In term of temperature, 80% of agents used temperature 0.7, while the rest used 0.2 or 0.3. Input/output token patterns shows MMLU-Pro’s best agent used more tokens than the average (+13% input, +17% output), while the other tasks showed minimal token differences. No agent incorporated worked examples or few-shot demonstrations. Finally, the agent generation performance remains similar across datasets regardless of whether 1 or 5 task examples are provided, demonstrating that models can understand new tasks from a single example.

3 DATA-DRIVEN AGENT IMPROVEMENT

We next explore whether LLMs can not only create agents but also improve them over time through iterative optimization, using meta-agents to analyze performance and generate better configurations.

We implement a simple self-improvement orchestration pattern involving four agents: a main task agent A_{task} that is the target of improvement, evaluator A_{eval} , analyzer A_{analyze} , and improver A_{improve} . We employ simple scorers as A_{eval} rather than LLM-based ones to ensure reliable performance measurements. Although these experiments focus on optimizing A_{task} , the analyzer and improver agents themselves can serve as optimization targets in the same iterative process.

More precisely, given a development dataset $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$ and performance metric ρ , we seek to find agents A'_{task} that improve $\sum_i \rho(A'_{\text{task}}(x_i), y_i)$.

Improvement Loop: At iteration t , the system executes:

$$\text{Out}_t = A_{\text{task}}^{t-1}(\mathcal{D}, \text{Ctx}) \quad (1)$$

$$\text{Eval}_t = A_{\text{eval}}(\text{Out}_t, \mathcal{D}) \quad (2)$$

$$A_t = A_{\text{analyze}}(\text{Eval}_t, \text{Out}_t, \mathcal{D}) \quad (3)$$

$$A_{\text{task}}^t = A_{\text{improve}}(A_{\text{task}}^{t-1}, A_{\text{task}}^{\text{best}}, A_t, \mathcal{H}_{t-1}, \text{Ctx}) \quad (4)$$

$$\mathcal{H}_t = \mathcal{H}_{t-1}; (\text{Out}_t, \text{Eval}_t, A_t) \quad (5)$$

where Out_t are task outputs, Eval_t contains evaluation metrics (aggregate and individual results), A_t provides analytical insights, and $\mathcal{H}_t$ is the improvement history.

The analyzer’s role is to identify patterns in performance and provide qualitative insights; we define an analyzer that uses aggregate results and randomly samples correct and incorrect examples from the data. The improver uses current and best task agents, current performance analysis and historical

data, and is instructed to write a new agent while respecting the *modifiable* flag. The process starts with an initial agent A_{task}^0 , and terminates when a target performance is met or maximum number of iterations is reached. With respect to the optimization objective, the loop implements metric-agnostic improvement by embedding the target metric within the evaluator, allowing the meta-agents to implicitly understand and optimize toward the desired goal.

The Analyzer and Improver agents are given in Appendix E and F respectively, while their prompt templates are summarized in Appendix I. Importantly, as it can be observed, these agents have minimal prompts and rely on the documentation artifacts that are provided to them (Ctx): the README associated to the LLM framework (as in the previous section) and an additional README file describing the iterative improvement loop.

The Analyzer and Improver agents are detailed in Appendices E and F, with prompt templates in Appendix I. These meta-agents use minimal prompts and instead rely on documentation (Ctx): the LLM framework README and an additional README describing the iterative improvement process. This is an artifact-driven approach where rich contextual information replaces engineered prompts.

3.1 EXPERIMENTS

We included math tasks in our experiments because they remain challenging for pre-reasoning state-of-the-art models and exhibit test-time scaling where longer outputs improve accuracy (Snell et al., 2024; Muennighoff et al., 2025). This property creates a natural optimization space: agents can achieve better quality through longer reasoning chains at higher computational cost, or maintain quality while reducing token usage, making them ideal testbeds for iterative improvement loops.

This section introduces additional pseudo-random number generator (PRNG) datasets specifically designed to test test-time scaling properties. Unlike other tasks where models perform pattern recognition, PRNG sequence prediction offers no shortcuts and requires precise algorithmic execution. We use three PRNG algorithms (BBS/LCG/MWC) with deliberately small constants to be more approachable for LLMs, which do not excel at large number arithmetic. We set the task of predicting the k -th number in the sequence with $k = 2$, based on the observation that even strong models perform poorly on this task. In contrast, reasoning models (GPT-OSS-120b) generate long reasoning traces and achieve 100% accuracy even for larger k values. We implement an initial PRNG solver agent as a generic CoT code execution agent. See Appendix G for details on data set creation and PRNG agent.

3.1.1 ACCURACY OPTIMIZATION

We conduct 10-iteration improvement cycles using the same Analyzer/Improver across all tasks. We test Sonnet 3.5 v2 and GPT-OSS-120b as meta-agent LLMs and restrict task models to Sonnet 3.5 v2. As in previous sections, we use an Evaluator that computes accuracy and provides aggregate and per-problem scores.

Table 3 shows results on development and test splits. The “Original agent” uses the standard prompt from the previous section and serves as the starting point for all improvement loops. The “Best agent” column shows the highest-performing configuration from the previous section’s agent generation experiments (performed on the development set).

AIME and MMLU-Pro All agents achieve improved performance on development sets. The challenging AIME datasets show modest improvements ($0.22 \rightarrow 0.27$ with Sonnet 3.5 v2). MMLU-Pro demonstrates strong overall performance, with substantial development improvements ($0.82 \rightarrow 0.89/0.90$) translating to more modest but consistent test gains. Overall, the results indicate that only larger development improvements transfer to test gains, suggesting insufficient development data or that baseline configurations had already undergone significant optimization in prior work. In terms of solutions found, the best AIME agent identifies problem types using word matching and provides targeted guidance such as “Count systematically by cases” for grid problems and “Factor completely” for number theory problems.

Appendix H shows a typical analyzer output from Sonnet analyzers. On AIME for example, analyzers commonly: identify where agents make arithmetic errors or use incorrect formulas, observe performance across problem types (e.g. low on geometry), highlight cases where multiple solu-

Table 3: Performance of initial configurations and after improvement loops. Original agent is the standard task prompt and serves as the A_{task}^0 . Best agent reports on the best configuration detected in the previous section.

Dataset	Original Agent	Best Agent	Improvement Loop	
			Sonnet 3.5 v2	GPT-OSS-120b
AIME2024 I (dev)	0.22	0.25	0.27	0.25
AIME2024 II	0.07	0.05	0.11	0.03
AIME2025	0.05	0.03	0.03	0.03
MMLU-Pro-dev	0.82	0.90	0.89	0.84
MMLU-Pro	0.79	0.83	0.80	0.77
PRNG-bbs	0.21	-	0.57	1.0
PRNG-lcg	0.0	-	0.02	1.0
PRNG-mwc	0.0	-	1.0	1.0

tions yield different answers (and recommend lower temperature), reason about the optimization trajectory, or highlight the need for more exploration by increasing temperature. Improvement loops (plotted in Appendix K) are not monotonic, reflecting that the optimization process is mostly exploratory, with agents often showing temporary performance degradation before discovering more effective approaches.

PRNG tasks For PRNG tasks, the two meta-agents discovered very different solutions. By iterations 3 or 4, GPT-OSS-120b invariantly discovered agents that pre-compute correct answers and instruct the LLM to output them. Specifically, the Analyzer, which sees correct/incorrect examples, inferred that all inputs list the same algorithm and instructed the Improver to create agents that execute this algorithm. The exact implementation varied: for example inserting a prompt section “verification_note” which instructs the LLM to check the answer against the pre-computed answer or a section “generate_example_trace” which despite the name contains the actual computed execution trace for that data point. While not technically cheating, these solution assume that the observed examples are representative of the true data distribution. In contrast the Sonnet loops only discovered this solution for the MWC algorithm and failed to find a working solution for the LCG algorithm. For BBS, Sonnet 3.5 finds improved solutions which list explicit modulo arithmetic rules (which the Analyzer identifies as difficult), add examples and clear validation guidance, but do not compute the function deterministically (reaching 57% performance). Regarding test-time scaling, the system did not identify it as a reliable performance improvement strategy. While one improvement loop produced agents that generated 50% more output tokens (with 10% longer input), this increased verbosity did not consistently translate to better performance. Appendix J shows an extreme case where detailed reasoning instructions resulted in outputs twice the length of the starting agent.

Variation across loops We observe that results vary when running identical improvement loops multiple times. To explore this, we run each Sonnet loop 4 additional times and observe max scores in the range [0.84-0.89] for MMLU-Pro, [0.24-0.28] for AIME, and [0.57-0.93] for PRNG-bbs. Combined with the observation that larger development improvements generalize better, this suggests the improvement loop in its current form serves as a discovery tool rather than monotonic optimization. Despite this variability, a *single* 10-iteration loop discovers superior solutions within an effectively infinite search space of pre-processing functions, parameters, and prompt templates, showing very good optimization efficiency.

3.1.2 ACCURACY+COST OPTIMIZATION

We next explore optimization under different objectives using more complex initial configurations. To balance accuracy and efficiency, we introduce a cost-accuracy evaluator where cost equals input_tokens + 3 × output_tokens (reflecting real-world pricing). The scoring function assigns 0 points for incorrect answers and 1/cost for correct ones, prioritizing accuracy while minimizing cost as a secondary objective. We modify optimization goals by changing the evaluator’s scoring computation and description. Figure 1 shows fragments of an evaluator output.

```

378 1  "overall_accuracy": 0.267,
379 2  "overall_score": 2.1e-05,
380 3  "avg_input_tokens": 13006.1,
381 4  ...
382 5  "optimization_goal": "Minimize cost while maintaining accuracy",
383 6  "scoring_explanation": "Score: 0 if wrong answer, 1/cost if [...]"

```

Figure 1: Example evaluator output for cost+accuracy optimization. The explanation fields are intended to guide the optimization goals of the meta-agents

We experiment on AIME and MMLU-Pro datasets. MMLU-Pro uses the CoT+ICL agent from the previous section as A_{task}^0 . For AIME experiments, we start with an ICL agent that samples 2 examples per problem from the slk dataset (Muennighoff et al., 2025), which contains similar problems with reasoning traces from powerful models. This dataset has been shown to improve AIME performance when used to fine-tune models.

Table 4: Performance of initial configurations and after improvement loops for cost+accuracy optimization: Accuracy and average input/output tokens.

Dataset	Original Config			Improvement Loop Sonnet 3.5 v2			Improvement Loop GPT-OSS-120b		
	Acc	In	Out	Acc	In	Out	Acc	In	Out
AIME2024 I (dev)	0.24	12.8k	340	0.33	6.1k	347	0.29	478	197
AIME2024 II	0.09	13.9k	368	0.08	5.9k	350	0.07	483	209
AIME2025	0.05	14.2k	365	0.05	6.0k	334	0.03	488	216
MMLU-Pro (dev)	0.82	760	190	0.84	301	176	0.84	230	196
MMLU-Pro	0.79	805	192	0.78	256	172	0.79	275	196

Results are shown in Table 4. Both Sonnet and GPT models explored much more diverse solutions than in previous experiments, likely due to the more complex starting agent (for AIME) and the ambiguous cost function allowing both accuracy and cost optimizations. For AIME, agents implemented strategies including: 1) automatic problem categorization into geometric, algebraic, and combinatorial types, 2) targeted prompts for each problem type (e.g., for geometry: “1. Draw and label key elements, 2. List given measurements and relationships...”), and 3) type-based ICL example selection. This approach yielded the winning Sonnet solution. However, the winning GPT-OSS-120b agent (Appendix L) used none of these techniques, instead replacing random ICL sampling with a fixed geometry example. While cost-effective, this did not generalize to good test performance. The variation in discovered solutions suggests the development set may be too small for reliable AIME optimization. On MMLU-Pro, all agents were able to find much more efficient solutions in terms of cost with roughly 70% reduction in input tokens and no impact on test performance.

Overall the cost+accuracy optimization showed the potential for creative solution discovery, with agents autonomously developing problem classification systems, targeted prompting strategies, and token compression techniques. More sophisticated improvement patterns could incorporate larger development sets, extended exploration phases, or multi-objective optimization strategies in order to further expand the range of discoverable solutions.

4 RELATED WORK

Large Language Models as Optimizers LLMs have emerged as powerful optimization tools across diverse domains. Automatic prompt engineering (Ramnath et al., 2025; Zhou et al., 2022; Li et al., 2025a; Debnath et al., 2025) has demonstrated LLMs’ ability to generate and refine prompts, with APE Zhou et al. (2022) introducing instructions as *programs* that can be systematically optimized. Similar work addresses automatic model selection (Wang et al., 2023; Tanaka et al., 2023; Tang et al., 2024), ICL example selection (Liu et al., 2022; Zhang et al., 2022; Do et al., 2024),

or prompt compression (Jiang et al., 2023; Li et al., 2025b) among many others. Recent advances explore LLMs as general-purpose optimizers (Jiang et al., 2025a; Yang et al., 2024a; Fernando et al., 2023), leveraging natural language understanding to interpret problems, generate solutions, and analyze feedback iteratively.

Building on these, our work introduces workflow optimization that extends beyond individual components to optimize entire agent configurations simultaneously: the system shows solution discovery capabilities that go beyond prompt variations. The approach enables metric-agnostic optimization through modular evaluator components, supporting flexible natural language-expressed optimization goals. Additionally, the self-referential optimization where meta-agents are defined within the same configurable framework as the task agents they optimize, enables the improvement process to enhance itself. Finally we designed this as an artifact-driven framework which leverages rich contextual documentation and examples to create a self-evolving system that accumulates knowledge rather than requiring extensive meta-prompt engineering.

Autonomous Agents and LLM Workflow Optimization Frameworks Research in LLM-based systems has advanced along two complementary but very related directions: autonomous agents for dynamic problem-solving and optimization frameworks for systematic workflow improvement.

Autonomous Agent Systems have demonstrated remarkable capabilities in dynamic environments. ReAct Yao et al. (2023) integrates reasoning with external tools through “thought-action-observation” loops, enabling LLMs to reason, perform actions, and learn from outcomes. PAL Gao et al. (2022) focuses on programmatic reasoning by generating executable code as intermediate steps, e.g. offloading complex calculations to interpreters for improved accuracy on mathematical tasks. AutoGPT and similar frameworks create agents that independently set goals, break down tasks, and execute multi-step plans with minimal human intervention. AIDE Jiang et al. (2025b) introduces tree-search strategies for systematic code and ML pipeline optimization, structuring solutions hierarchically for targeted improvements based on incremental changes.

Many optimization frameworks have emerged to systematically improve LLM workflows. Some examples are DSPy (Khatab et al., 2023), which enables “programming, not prompting” through modular pipeline definition and automatic prompt optimization. TextGrad (Huang et al., 2023) introduces “automatic differentiation via text,” using LLM-based feedback for iterative refinement within computational graphs. GPTSwarm (Zhuge et al., 2024) represents agents as optimizable computational graphs, unifying prompt engineering techniques through node-level optimization and edge-level orchestration. Teola Zheng et al. (2024) provides end-to-end optimization through dataflow graph representations, enabling rule-based optimizations across workflow modules.

Our approach bridges these paradigms through simple primitives that enable structured, interpretable optimization that combines the systematic exploration of optimization frameworks with the adaptive improvement mechanisms of autonomous agents. Unlike open-ended autonomous exploration that can produce black-box improvements, the configuration-driven approach ensures all generated solutions remain human-verifiable. Instead of replacing humans, the system can serve as a solution discovery tool that aids human decision-making through transparent improvement trajectories. This is possible through the constrained but expressive solution space, resulting in a middle ground between rigid optimization frameworks and unconstrained exploration systems.

5 FUTURE WORK

Current improvement cycles analyze development data point-by-point. While necessary for runtime execution, development stages could benefit from global dataset processing to enable aggregate analysis and dataset-wide optimization strategies not visible in individual examples. Additionally, more sophisticated improvement patterns could incorporate multi-objective optimization or Monte Carlo solution search, and beyond fixed meta-agents, the configuration-driven design provides a foundation for joint optimization of both task and meta-agents.

Finally, the framework can be extended to support cross-domain knowledge transfer by accumulating improvement artifacts across domains. This would build a repository of successful patterns and strategies, with accumulated knowledge informing improvements on new tasks and enabling transfer learning at the agent architecture level.

REFERENCES

- Zina Chkirbene, Ridha Hamila, Ala Gouisse, and Unal Devrim. Large language models (llm) in industry: A survey of applications, challenges, and trends. In *2024 IEEE 21st International Conference on Smart Communities: Improving Quality of Life using AI, Robotics and IoT (HONET)*, pp. 229–234, 2024. doi: 10.1109/HONET63146.2024.10822885.
- Tonmoy Debnath, Nurul Absar Siddiky, Muhammad Enayetun Rahman, Prosenjit Das, and Antu Kumar Guha. A comprehensive survey of prompt engineering techniques in large language models. *arXiv preprint arXiv:2503.04403*, 2025.
- Xuan Long Do, Yiran Zhao, Hannah Brown, Yuxi Xie, James Xu Zhao, Nancy F. Chen, Kenji Kawaguchi, Michael Shieh, and Junxian He. Prompt optimization via adversarial in-context learning. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 7308–7327, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.395. URL <https://aclanthology.org/2024.acl-long.395/>.
- Chrisantha Fernando, Dylan Banarse, Henryk Michalewski, Simon Osindero, and Tim Rocktäschel. Promptbreeder: Self-referential self-improvement via prompt evolution, 2023. URL <https://arxiv.org/abs/2309.16797>.
- Luyu Gao, Aniruddha Madaan, Seyed Shakeri, Diyi Yu, Ling Wang, Jun Han, Han Chen, and Jian Lin. PAL: Program-aided language models. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pp. 577–588, 2022.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring mathematical problem solving with the MATH dataset. In *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1*, NeurIPS Datasets and Benchmarks 2021, virtual, December 2021. URL <https://datasets-benchmarks-proceedings.neurips.cc/paper/2021/hash/be83ab3ecd0db773eb2dc1b0a17836a1-Abstract-round2.html>.
- Tianyi Huang, Huimin Yao, Jiacheng He, Zhiyi Zheng, and Zhou Lin. TextGrad: Autodifferentiating through arbitrary text with large language models. *arXiv preprint arXiv:2308.02450*, 2023.
- Caigao Jiang, Xiang Shu, Hong Qian, Xingyu Lu, Jun Zhou, Aimin Zhou, and Yang Yu. Llmopt: Learning to define and solve general optimization problems from scratch, 2025a. URL <https://arxiv.org/abs/2410.13213>.
- Huiqiang Jiang, Qianhui Wu, Chin-Yew Lin, Yuqing Yang, and Lili Qiu. Llmilingua: Compressing prompts for accelerated inference of large language models, 2023. URL <https://arxiv.org/abs/2310.05736>.
- Zhengyao Jiang, Dominik Schmidt, Dhruv Srikanth, Dixing Xu, Ian Kaplan, Deniss Jacenko, and Yuxiang Wu. Aide: Ai-driven exploration in the space of code. 2025b. URL <https://arxiv.org/abs/2502.13138>.
- Omar Khattab, Arnav Singh, Zhiqing Li, Hamza Khalid, Yi Zhang, Mirian Sanjabi, Amir Gholami, Matei Zaharia, and Matei Zaharia. DSPy: Compiling declarative language model calls into state-of-the-art pipelines. *arXiv preprint arXiv:2311.00062*, 2023.
- Jiaheng Li, Junhao Hu, Jionghao Xu, Yuxin Liu, Yilun Zhang, Li Huang, Ruichen Zhao, Yunpu Wang, Jihong Liu, and Meng Wu. Large language models for constructing and optimizing machine learning workflows: A survey. *arXiv preprint arXiv:2411.10478*, 2024.
- Wenwu Li, Xiangfeng Wang, Wenhao Li, and Bo Jin. A survey of automatic prompt engineering: An optimization perspective. *arXiv preprint arXiv:2502.11560*, 2025a.
- Zongqian Li, Yinhong Liu, Yixuan Su, and Nigel Collier. Prompt compression for large language models: A survey. In Luis Chiruzzo, Alan Ritter, and Lu Wang (eds.), *Proceedings of the 2025*

540 *Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 7182–7195, Albuquerque,
541 New Mexico, April 2025b. Association for Computational Linguistics. ISBN 979-8-89176-189-
542 6. doi: 10.18653/v1/2025.naacl-long.368. URL <https://aclanthology.org/2025.naacl-long.368/>.
543
544
545 Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan
546 Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. In *The Twelfth
547 International Conference on Learning Representations*, Vienna, Austria, May 7-11 2024. Open-
548 Review.net. URL <https://openreview.net/forum?id=v8L0pN6EOi>.
549
550 Jiachang Liu, Dinghan Shen, Yizhe Zhang, Bill Dolan, Lawrence Carin, and Weizhu Chen. What
551 makes good in-context examples for GPT-3? In Eneko Agirre, Marianna Apidianaki, and Ivan
552 Vulić (eds.), *Proceedings of Deep Learning Inside Out (DeeLIO 2022): The 3rd Workshop on
553 Knowledge Extraction and Integration for Deep Learning Architectures*, pp. 100–114, Dublin,
554 Ireland and Online, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/
555 2022.deelio-1.10. URL <https://aclanthology.org/2022.deelio-1.10/>.
556
557 MAA Committees. AIME Problems and Solutions. Art of Problem Solving Wiki. URL
558 [https://artofproblemsolving.com/wiki/index.php/AIME_Problems_](https://artofproblemsolving.com/wiki/index.php/AIME_Problems_and_Solutions)
559 [and_Solutions](https://artofproblemsolving.com/wiki/index.php/AIME_Problems_and_Solutions). Accessed: September 25, 2025.
560
561 Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke
562 Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. s1: Simple test-time
563 scaling, 2025. URL <https://arxiv.org/abs/2501.19393>.
564
565 Kiran Ramnath, Kang Zhou, Sheng Guan, Soumya Smruti Mishra, Xuan Qi, Zhengyuan Shen, Shuai
566 Wang, Sangmin Woo, Sullam Jeoung, Yawei Wang, Haozhu Wang, Han Ding, Yuzhe Lu, Zhichao
567 Xu, Yun Zhou, Balasubramaniam Srinivasan, Qiaojing Yan, Yueyan Chen, Haibo Ding, Panpan
568 Xu, and Lin Lee Cheong. A systematic survey of automatic prompt optimization techniques,
569 2025. URL <https://arxiv.org/abs/2502.16923>.
570
571 Mubashar Raza, Zarmina Jahangir, Muhammad Riaz, and Muhammad Sattar. Industrial applications
572 of large language models. *Scientific Reports*, 15, 04 2025. doi: 10.1038/s41598-025-98483-1.
573
574 Yasmeen Saleh, Manar Abu Talib, Qassim Nasir, and Fatima Dakalbab. Evaluating large language
575 models: A systematic review of efficiency, applications, and future directions. *Frontiers in Com-
576 puter Science*, 7:1523699, 2025. doi: 10.3389/fcomp.2025.1523699.
577
578 Patrick Schramowski, Max Eichenseer, Konstantina Stamati, Christian Klein, David Roth,
579 and Christian Bizer. Exploring the role of large language models in the scientific pro-
580 cess. *Nature Machine Intelligence*, 2025. URL [https://www.nature.com/articles/
581 s44387-025-00019-5](https://www.nature.com/articles/s44387-025-00019-5).
582
583 Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling llm test-time compute optimally
584 can be more effective than scaling model parameters, 2024. URL [https://arxiv.org/
585 abs/2408.03314](https://arxiv.org/abs/2408.03314).
586
587 Shun Tanaka, Takehiko Kadowaki, Yuya Nagai, Yuki Nishiyama, and Masaharu Kakiuchi. Autonlp:
588 A framework for automated model selection in natural language processing. In *2023 IEEE 47th
589 Annual Computers, Software, and Applications Conference (COMPSAC)*, pp. 2011–2016. IEEE,
590 2023.
591
592 Jiabin Tang, Lianghao Xia, Zhonghang Li, and Chao Huang. Ai-researcher: Autonomous scientific
593 innovation, 2025. URL <https://arxiv.org/abs/2505.18705>.
594
595 Shicheng Tang, Yongxu Xu, Xiang Li, Zhixu Shi, and Chengxi Zhai. Automatic prompt selection
596 for large language models. *arXiv preprint arXiv:2404.02717*, 2024.
597
598 Shizhe Wang, Xiaoyu Hu, Yu He, Guoliang Xu, Hongguang Liu, and Zhenhua Lin. Automatic
599 model selection with large language models for mathematical reasoning. In *Findings of the Asso-
600 ciation for Computational Linguistics: EMNLP 2023*, pp. 835–845, 2023.

- Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni, Abhuranil Chandra, Shiguang Guo, Weiming Ren, Aaran Arulraj, Xuan He, Ziyang Jiang, Tianle Li, Max Ku, Kai Wang, Alex Zhuang, Rongqi Fan, Xiang Yue, and Wenhui Chen. Mmlu-pro: A more robust and challenging multi-task language understanding benchmark, 2024. URL <https://arxiv.org/abs/2406.01574>.
- Ziyue Wang and et al. Ai-enabled scientific revolution in the age of generative ai. *Nature*, 1(1): 1–10, 2025. doi: 10.1038/s44387-025-00018-6.
- Yuanzhi Xiao, Guoxin Li, Ziyue Wang, Xinyu Zhao, Zhenhua Lin, Chenxiao Liu, Fan Chen, Chen Zhang, Guoping Liu, and Qi Li. Prompt cache: Modular attention reuse for low-latency inference. *arXiv preprint arXiv:2311.04934*, 2023.
- Chengrun Yang, Xuezhi Wang, Yifeng Lu, Hanxiao Liu, Quoc V Le, Denny Zhou, and Xinyun Chen. Large language models as optimizers. In *The Twelfth International Conference on Learning Representations*, 2024a. URL <https://openreview.net/forum?id=Bb4VGOWELI>.
- Tianxing Yang, Yilun Zhang, Jiaheng Li, Junhao Hu, and Jionghao Xu. Hypothesis generation with large language models. *arXiv preprint arXiv:2404.04326*, 2024b.
- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Luo, Kai Xu, Jian Li, Jian Li, George Liang, Tianyu Huang, and Bo Hu. ReAct: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2302.05380*, 2023.
- Yiming Zhang, Shi Feng, and Chenhao Tan. Active example selection for in-context learning. In Yoav Goldberg, Zornitsa Kozareva, and Yue Zhang (eds.), *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pp. 9134–9148, Abu Dhabi, United Arab Emirates, December 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.emnlp-main.622. URL <https://aclanthology.org/2022.emnlp-main.622/>.
- Junyi Zheng, Ziyue Wang, Xinyu Zhao, Guoxin Li, Zekun Wang, Zhenhua Lin, Chenxiao Liu, Fan Chen, Chen Zhang, Guoping Liu, and Qi Li. Teola: Towards end-to-end optimization of llm-based applications. *arXiv preprint arXiv:2407.00326*, 2024.
- Yongchao Zhou, Andrei Song, Shizong Ma, Jiao Zhang, Yuhang Han, Tianci Liu, Ning Ma, Quanquan Liu, Guang Sun, Ming Xu, Hong Yang, Ruibin Li, and Zhijie Zeng. Large language models are human-level prompt engineers. *arXiv preprint arXiv:2211.01910*, 2022.
- Banghua Zhu, Ying Sheng, Lianmin Zheng, Clark Barrett, Michael I. Jordan, and Jiantao Jiao. On optimal caching and model multiplexing for large model inference, 2023. URL <https://arxiv.org/abs/2306.02003>.
- Mingchen Zhuge, Wenyi Wang, Louis Kirsch, Francesco Faccio, Dmitrii Khizbullin, and Jürgen Schmidhuber. Gptswarm: language agents as optimizable graphs. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org, 2024.

A GENERATIVE AI USE

AI tools were used throughout this research to enhance productivity and improve the quality of the work.

Framework Development Process The framework and the framework documentation was developed iteratively with assistance from a code assistant powered by Sonnet 3.5 and Sonnet 4.0, depending on availability. The assistant was given specific instructions to improve documentation clarity and update the framework based on feedback. This iterative process was validated by testing the framework on toy prediction tasks (separate from the experimental tasks) to ensure the generated agents functioned correctly.

648 **Paper writing** LLMs assisted in refining the clarity and coherence of the writing, suggesting more
649 concise phrasings and identifying areas where technical concepts could be explained more clearly
650 for broader accessibility. Additionally, AI tools facilitated data analysis and presentation by generat-
651 ing LaTeX tables and figures from raw experimental results. All AI-generated content was carefully
652 reviewed and validated by the authors to ensure accuracy and appropriateness for the scientific con-
653 text.

B LLM AGENT IMPLEMENTATION

LLM Agent Implementation At the level of implementation, we describe an LLM agent via three components: Configuration (json), Agent Class Implementation (Python) and Input Schema File.

The **configuration** file defines the agent’s behavior, model settings, input processing, prompt structure, and output format. A configuration tuple C consists of:

$$C = \langle \text{desc}, \text{schema}, \text{model}, \text{inputs}, \text{prompt}, \text{output}, \text{settings} \rangle$$

where:

- desc: Human-readable description of agent purpose
- input schema: JSON Schema path defining expected input data structure (including additional context data).
- model = $\langle \text{name}, \rho \rangle$: LLM specification (θ_{LLM} and inference parameters ρ)
- inputs = $\{f_1, \dots, f_k\}$: Data transformation functions operating on x_i and context c
- prompt = $\langle T, \text{sections} \rangle$: Template T with reusable sections and placeholders
- output = $\langle \text{format}, \text{schema}, \text{errors} \rangle$: Expected output format which drives the LLM response parsing
- settings: Execution parameters such as timeouts

Notably, inputs can be of two types: “data” or “computed” (with function name and arguments), in which case the function needs to be defined in the agent implementation class. For example the following configuration fragment shows the input section and references *raw_problem_data* which needs to exist in the input data processed, and *problem_text* which is computed via a function and takes the raw data as input.

Listing 1: Configuration example showing computed inputs

```
1 "inputs": {
2   "raw_problem_data": {
3     "source": "data"
4   }
5   "problem_text": {
6     "source": "computed",
7     "function": "extract_problem_text",
8     "args": {"problem_data": "raw_problem_data"}
9   }
10 }
```

The role of the **agent class** is to implement the functions referenced in the configuration’s computed inputs. These functions operate at data-point level and can apply custom pre-processing (such as normalization, feature extraction, problem type classification, etc), retrieve examples from static resources provided as context c , etc. In the example above, the python agent class needs to implement the *extract_problem_text* function, for example as:

```
1 def extract_problem_text(self, problem_data):
2   """Function referenced in config's computed inputs"""
3   return problem_data.get('problem', '')
```

The **input schema file** and the output section of the configuration define the communication with upstream and downstream tasks. The agent can chose to ignore parts of the input data but can’t access any resources outside what is defined in the input schema file.

Modifiability Control In order to simplify the process of other LLMs improving task agents, we introduce modifiability control. Each configuration section includes a flag *modifiable* $\in \{\text{true}, \text{false}\}$ that guides other agents wrt which sections can be modified. For example, *model.modifiable* = true indicates that model name or model hyper-parameters can be modified, while *output.modifiable* = false indicates that the response schema needs to be preserved. This enables systematic exploration

756 of the configuration space while maintaining structural invariants essential for downstream process-
757 ing, as well as the option to configure preferences (for example to prevent changes in the underlying
758 LLM model if so desired).

759 **Agent Execution** The agents operate synchronously, processing one input data point at a time in
760 sequential order. The framework reads the agent definition from three files (JSON configuration,
761 Python implementation, and input schema), performs static validation of the configuration and dy-
762 namic validation of runtime input data. At runtime, the framework processes input data through
763 defined transformations, generates the final prompt by resolving placeholders, and calls the LLM
764 with the rendered prompt. The LLM response is then parsed by an output parser according to the
765 output format specified in the configuration.
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809

C MATH PROBLEM SOLVING AGENT

Listing 2: AIME 2024 Task Agent Configuration

```
1 {
2   "description": "AIME 2024 task agent - generates detailed reasoning followed by 3-digit
3     integer answer",
4   "input_schema": "agents/task/task_input_schema.json",
5   "model": {
6     "_modifiable": true,
7     "name": "anthropic.claude-3-5-sonnet-20241022-v2:0",
8     "parameters": {
9       "temperature": 0.7,
10      "max_tokens": 8192
11    }
12  },
13  "inputs": {
14    "_modifiable": true,
15    "problem_text": {
16      "source": "computed",
17      "function": "extract_problem_text",
18      "args": {"problem_data": "raw_problem_data"},
19      "description": "Extracted problem text from AIME2024 format"
20    },
21    "raw_problem_data": {
22      "source": "data",
23      "path": "$",
24      "required": true,
25      "description": "Complete problem data from AIME2024 dataset"
26    }
27  },
28  "prompt": {
29    "_modifiable": true,
30    "sections": {
31      "instructions": "You are a mathematical problem solver specializing in AIME problems.",
32      "task_description": "AIME problems require integer answers between 000 and 999. Show
33        detailed step-by-step reasoning, then provide the final numerical answer.",
34      "output_format": "You must respond with a JSON object containing exactly two fields:\n-
35        \"reasoning\": string containing your detailed step-by-step solution\n- \"answer\":
36        string containing your 3-digit answer (pad with leading zeros if needed)"
37    },
38    "template": "### Instructions:\n{prompt.sections.instructions}\n\n### Task:\n{prompt.
39      sections.task_description}\n\n### Problem:\n{inputs.problem_text}\n\n### Output
40      Format:\n{prompt.sections.output_format}\n\nNow solve the problem and respond with
41      the JSON object:"
42  },
43  "output": {
44    "_modifiable": false,
45    "format": "json",
46    "schema": {
47      "type": "object",
48      "required": ["reasoning", "answer"],
49      "properties": {
50        "reasoning": {"type": "string", "description": "Detailed step-by-step solution"},
51        "answer": {"type": "string", "pattern": "[0-9]{3}", "description": "3-digit integer
52          answer (000-999)"}
53      }
54    },
55    "error_values": {
56      "reasoning": "Error occurred during reasoning.",
57      "answer": "PARSE_ERROR"
58    }
59  },
60  "settings": {
61    "_modifiable": false,
62    "class_name": "AIME2024TaskAgent",
63    "timeout_seconds": 300
64  }
65 }
```

Listing 3: AIME 2024 Task Agent Implementation

```
1 from llm_agent import LLMAgent
2 from typing import Dict, Any
3
4 class AIME2024TaskAgent(LLMAgent):
5     """
6     AIME2024-specific reasoning + answer agent.
7     """
```

```

864 8      Processes AIME mathematical competition problems and generates detailed step-by-step
865 9      reasoning followed by answers.
866 10     """
867 11     def __init__(self, name="AIME2024_Task", config_path=None, dry_run=False):
868 12         super().__init__(name, config_path, dry_run)
869 13
870 14     def extract_problem_text(self, problem_data: Dict[str, Any]) -> str:
871 15         """
872 16         Extract problem text from AIME2024 data format.
873 17
874 18         AIME2024 format:
875 19         {
876 20             "id": int,
877 21             "problem": str, # Main problem statement
878 22             "url": str
879 23         }
880 24         """
881 25         return str(problem_data.get('problem', ''))

```

Listing 4: AIME 2024 Task Agent Input Schema

```

880 1  {
881 2  "type": "object",
882 3  "required": ["id", "problem"],
883 4  "properties": {
884 5      "id": {
885 6          "type": "integer",
886 7          "description": "Problem identifier"
887 8      },
888 9      "problem": {
889 10         "type": "string",
890 11         "description": "The AIME problem statement"
891 12     },
892 13     "url": {
893 14         "type": "string",
894 15         "description": "URL to problem source (optional)"
895 16     }
896 17 }
897 18 }

```

D FRAMEWORK CAPABILITIES

The LLM Framework provides extensive flexibility for implementing diverse agent architectures and optimization strategies through its configuration-driven design. The framework's expressive power enables sophisticated agent implementations across multiple dimensions:

Feature Extraction and ML Integration: Agents can incorporate external machine learning models and feature extraction pipelines through initialization methods. Complex preprocessing can be implemented in the agent's `--init--` method, loading pre-trained models, statistical analyzers, or domain-specific feature extractors. External datasets for feature computation can be specified in the configuration's input section and loaded during agent initialization:

```
"inputs": {
  "feature_dataset": {
    "source": "data",
    "path": "external_features.json",
    "description": "External dataset for feature extraction"
  },
  "extracted_features": {
    "source": "computed",
    "function": "extract_ml_features",
    "args": {"raw_data": "problem_data", "feature_db": "feature_dataset"}
  }
}
```

In-Context Learning Implementation: ICL can be most flexibly implemented as computed inputs that dynamically generate examples using various example selection strategies, including similarity-based retrieval, difficulty-matched sampling, or domain-specific filtering. ICL examples can be sourced from multiple datasets with automatic data leakage prevention:

```
"icl_examples": {
  "source": "computed",
  "function": "generate_icl_examples",
  "args": {
    "icl_data_path": "data/training_examples.json",
    "num_examples": 5,
    "similarity_metric": "cosine",
    "exclude_current": true
  }
}
```

Dynamic Input Transformation: The computed input system enables arbitrary feature engineering and input pre-processing. Agents can implement domain-specific transformations, adaptive input formatting, etc. Each transformation function can access multiple input sources and apply complex processing logic:

```
"processed_input": {
  "source": "computed",
  "function": "transform_input",
  "args": {
    "text_data": "raw_text",
    "numerical_data": "raw_numbers",
    "context_data": "external_context",
    "transformation_type": "hierarchical_fusion"
  }
}
```

Model and Parameter Optimization: The framework supports dynamic model selection and parameter tuning through the modifiable model configuration. Agents can be optimized across dif-

972 ferent LLM models with varying capabilities and costs. Model parameters including temperature,
973 max_tokens, and model-specific settings can be systematically explored:
974

```
975 "model": {  
976   "_modifiable": true,  
977   "name": "anthropic.claude-3-5-sonnet-20241022-v2:0",  
978   "parameters": {  
979     "temperature": 0.7,  
980     "max_tokens": 4096,  
981     "top_p": 0.9  
982   }  
983 }
```

984 **Prompt Engineering and Structure:** The modular prompt system enables sophisticated prompt ar-
985 chitectures with independently optimizable sections. Agents can implement hierarchical prompting,
986 conditional prompt sections, dynamic example integration, and task-specific reasoning templates.
987 The template system supports complex placeholder substitution and structured prompt composition.

988 **Orchestration Constraints:** While the framework provides extensive flexibility, certain constraints
989 ensure orchestration compatibility. The output schema should remain fixed (`_modifiable:`
990 `false`) to maintain consistency across the evaluation-analysis-improvement loops. This constraint
991 ensures that orchestration agents can reliably process results while allowing freedom in input pro-
992 cessing, model selection, and reasoning strategies.
993

994
995
996
997
998
999
1000
1001
1002
1003
1004
1005
1006
1007
1008
1009
1010
1011
1012
1013
1014
1015
1016
1017
1018
1019
1020
1021
1022
1023
1024
1025

E ANALYZER AGENT

Listing 5: Analyzer Agent Configuration

```
1030 1 {
1031 2   "description": "Generic analyzer agent - analyzes task performance and identifies
1032 3   improvement opportunities",
1033 4   "input_schema": "analyzer_input_schema.json",
1034 5   "model": {
1035 6     "_modifiable": true,
1036 7     "name": "anthropic.claude-3-5-sonnet-20241022-v2:0",
1037 8     "parameters": {
1038 9       "temperature": 0.3,
1039 10      "max_tokens": 8192
1040 11    }
1041 12 },
1042 13 "inputs": {
1043 14   "_modifiable": true,
1044 15   "current_task_config": {
1045 16     "source": "data",
1046 17     "path": "current_task_config",
1047 18     "required": true,
1048 19     "description": "Current task agent configuration as JSON string"
1049 20   },
1050 21   "evaluation_summary": {
1051 22     "source": "data",
1052 23     "path": "evaluation_summary",
1053 24     "required": true,
1054 25     "description": "Aggregated evaluation metrics from evaluator"
1055 26   },
1056 27   "individual_results": {
1057 28     "source": "data",
1058 29     "path": "individual_results",
1059 30     "required": true,
1060 31     "description": "Per-problem results with model outputs, ground truth, and scores"
1061 32   },
1062 33   "correct_examples": {
1063 34     "source": "computed",
1064 35     "function": "extract_correct_examples",
1065 36     "args": {
1066 37       "individual_results": "individual_results",
1067 38       "num_examples": 3,
1068 39       "score_field": "score_field",
1069 40       "perfect_score": "perfect_score"
1070 41     },
1071 42     "description": "Examples where the task agent performed correctly"
1072 43   },
1073 44   "incorrect_examples": {
1074 45     "source": "computed",
1075 46     "function": "extract_incorrect_examples",
1076 47     "args": {
1077 48       "individual_results": "individual_results",
1078 49       "num_examples": 5,
1079 50       "score_field": "score_field",
1080 51       "perfect_score": "perfect_score"
1081 52     },
1082 53     "description": "Examples where the task agent made errors"
1083 54   }
1084 55 },
1085 56 "prompt": {
1086 57   "template": "### Instructions:\n{prompt.sections.instructions}\n\n### Current Task Agent
1087 58   Configuration:\n{inputs.current_task_config}\n\n### Evaluation Summary:\n{inputs.
1088 59   evaluation_summary}\n\n### Examples Where Task Agent Was CORRECT:\n{inputs.
1089 60   correct_examples}\n\n### Examples Where Task Agent Made ERRORS:\n{inputs.
1090 61   incorrect_examples}\n\nAnalyze the task agent performance and identify improvement
1091 62   opportunities:"
1092 63 },
1093 64 "output": {
1094 65   "_modifiable": false,
1095 66   "format": "json",
1096 67   "schema": {
1097 68     "type": "object",
1098 69     "required": ["analysis"],
1099 70     "properties": {
1100 71       "analysis": {
1101 72         "type": "string",
1102 73         "description": "Complete analysis of task agent performance"
1103 74       }
1104 75     }
1105 76   }
1106 77 }
```

```

1080 71 },
1081 72 "settings": {
1082 73     "_modifiable": false,
1083 74     "class_name": "AnalyzerAgent",
1084 75     "timeout_seconds": 600
1085 76 }
1086 77 }

```

Listing 6: Analyzer Agent Implementation

```

1088 1 from llm_agent import LLMAgent
1089 2 from typing import Dict, List, Any
1090 3 import json
1091 4 import random
1092 5
1093 6 class AnalyzerAgent(LLMAgent):
1094 7     """
1095 8     Generic analyzer agent for performance analysis.
1096 9
1097 10     Analyzes task agent performance and identifies improvement opportunities
1098 11     for any domain that follows the standard evaluator contract.
1099 12     """
1100 13
1101 14     def __init__(self, name="Analyzer", config_path=None, dry_run=False):
1102 15         super().__init__(name, config_path, dry_run)
1103 16
1104 17     def extract_correct_examples(self, individual_results: List[Dict[str, Any]],
1105 18                                num_examples: int,
1106 19                                score_field: str = "score",
1107 20                                perfect_score: float = 1.0) -> str:
1108 21         """
1109 22         Extract examples where the task agent performed correctly.
1110 23         """
1111 24         correct_examples = [result for result in individual_results
1112 25                             if result.get(score_field, 0) >= perfect_score]
1113 26
1114 27         selected = random.sample(correct_examples, min(num_examples, len(correct_examples)))
1115 28         return json.dumps(selected, indent=2)
1116 29
1117 30     def extract_incorrect_examples(self, individual_results: List[Dict[str, Any]],
1118 31                                  num_examples: int,
1119 32                                  score_field: str = "score",
1120 33                                  perfect_score: float = 1.0) -> str:
1121 34         """
1122 35         Extract examples where the task agent made errors.
1123 36         """
1124 37         incorrect_examples = [result for result in individual_results
1125 38                               if result.get(score_field, 0) < perfect_score]
1126 39
1127 40         selected = random.sample(incorrect_examples, min(num_examples, len(incorrect_examples)))
1128 41         return json.dumps(selected, indent=2)

```

F IMPROVER AGENT

Listing 7: Improver Agent Configuration

```
1138 1 {
1139 2   "description": "Generic improver agent - generates improved task agent based on analysis",
1140 3   "input_schema": "improver_input_schema.json",
1141 4   "model": {
1142 5     "_modifiable": true,
1143 6     "name": "anthropic.claude-3-5-sonnet-20241022-v2:0",
1144 7     "parameters": {
1145 8       "temperature": 0.3,
1146 9       "max_tokens": 25000
1147 10    }
1148 11  },
1149 12  "inputs": {
1150 13    "_modifiable": true,
1151 14    "analysis_results": {
1152 15      "source": "data",
1153 16      "path": "analysis_results",
1154 17      "required": true,
1155 18      "description": "Analysis results from the analyzer agent"
1156 19    },
1157 20    "current_task_agent": {
1158 21      "source": "data",
1159 22      "path": "current_task_agent",
1160 23      "required": true,
1161 24      "description": "Current task agent configuration"
1162 25    },
1163 26    "best_task_agent": {
1164 27      "source": "data",
1165 28      "path": "best_task_agent",
1166 29      "required": false,
1167 30      "description": "Best performing task agent configuration and implementation"
1168 31    }
1169 32  },
1170 33  "prompt": {
1171 34    "template": "### Instructions:\n{prompt.sections.instructions}\n\n### Current Task Agent:\n\n{inputs.current_task_agent}\n\n### Best Performing Agent:\n{inputs.best_task_agent}\n\n### Analysis Results:\n{inputs.analysis_results}\n\nGenerate improved task agent configuration based on the analysis:"
1172 35  },
1173 36  "output": {
1174 37    "_modifiable": false,
1175 38    "format": "json",
1176 39    "schema": {
1177 40      "type": "object",
1178 41      "required": ["explanation", "new_task_implementation", "new_task_config"],
1179 42      "properties": {
1180 43        "explanation": {
1181 44          "type": "string",
1182 45          "description": "Detailed explanation of improvements"
1183 46        },
1184 47        "new_task_implementation": {
1185 48          "type": "string",
1186 49          "description": "Complete Python file content"
1187 50        },
1188 51        "new_task_config": {
1189 52          "type": "string",
1190 53          "description": "Complete JSON config content"
1191 54        }
1192 55      }
1193 56    }
1194 57  },
1195 58  "settings": {
1196 59    "_modifiable": false,
1197 60    "class_name": "ImproverAgent",
1198 61    "timeout_seconds": 1200
1199 62  }
1200 63 }
```

Listing 8: Improver Agent Implementation

```
1184 1 from llm_agent import LLMAgent
1185 2
1186 3 class ImproverAgent(LLMAgent):
1187 4     """
1188 5     Generic improver agent - generates improved task agents based on analysis.
1189 6     """
```

```
1188 7
1189 8     def __init__(self, name="Improver", config_path=None, dry_run=False):
1190 9         super().__init__(name, config_path, dry_run)
```

```
1191
1192
1193
1194
1195
1196
1197
1198
1199
1200
1201
1202
1203
1204
1205
1206
1207
1208
1209
1210
1211
1212
1213
1214
1215
1216
1217
1218
1219
1220
1221
1222
1223
1224
1225
1226
1227
1228
1229
1230
1231
1232
1233
1234
1235
1236
1237
1238
1239
1240
1241
```

G PRNG EXPERIMENTS

We use the three algorithms BBS/LCG/MWC, and deliberately set small constants to be more approachable for LLMs which do not excel at large number arithmetic:

- **LCG:** $X_{n+1} = (25173 \cdot X_n + 13849) \bmod 65536$ - smaller modulus than typical implementations
- **MWC:** $X_{n+1} = (18782 \cdot X_n + \text{carry}) \& 0\text{xFFF}$ - 12-bit output with carry propagation
- **BBS:** $X_{n+1} = X_n^2 \bmod 209$ where $209 = 11 \times 19$ - small primes

We create train and test data sets (each containing 70 sequences) where models must predict the k 'th number in the sequence given: (1) complete algorithm implementation in Python, (2) hyperparameters (e.g. the prime numbers used) and (3) initial seed value (distinct in train/test splits). We set the task as that of predicting the second number in the sequence ($k = 2$) based on the initial observation that even state-of-the-art models perform poorly on the task. In contrast, we confirm that a reasoning model (OpenAI 120b) generates long reasoning traces and achieves 100% accuracy.

Table 5: PRNG agents for BBS/LCG/MWC algorithms. The names of the algorithms are obfuscated in case the LLMs are aware of these algorithms and have memorized sequences. `<Function>` is the Python code, `<N>` is the number to be predicted (in this case at position 2), `<Initial Value>` the seed, `<Hyperparams>` are constants used in the algorithms. Note that this is a generic code execution agent, as all the placeholders are only resolved at run-time when individual data points are processed.

Instructions: You are a function evaluator. You will be given a function implementation, its hyperparameters, and an initial value. Your task is to predict what the function will return at a specific position. Task: Given the function implementation and parameters, determine what value the function would return at position `<N>` when called with the initial value. Note: The first number generated is at position 1, the first iteration is position 1. Show your detailed step-by-step reasoning, then provide the final prediction.
Function Implementation: `<Function>`
Hyperparameters: `<Hyperparams>`
Initial Value: `<Initial Value>`
Approach: 1. Carefully analyze the function implementation 2. Understand what the hyperparameters control 3. Trace through the function execution step by step 4. For modulo operations (%), decompose them: compute the division first, then find the remainder - do NOT compute modulo in a single step 5. For ANY arithmetic expression with more than 2 terms (e.g., $a * b + c$, $x + y - z$, etc.), break it down step by step and compute intermediate results - do NOT evaluate complex expressions in one shot 6. Calculate the value at position `<N>` 7. Present your reasoning followed by the final answer
Output Format: JSON object with "reasoning" (detailed step-by-step analysis) and "answer" (predicted number as string) fields.

1296
1297
1298
1299
1300
1301
1302
1303
1304
1305
1306
1307
1308
1309
1310
1311
1312
1313
1314
1315
1316
1317
1318
1319
1320
1321
1322
1323
1324
1325
1326
1327
1328
1329
1330
1331
1332
1333
1334
1335
1336
1337
1338
1339
1340
1341
1342
1343
1344
1345
1346
1347
1348
1349

H ANALYZER OUTPUT EXAMPLES

Analysis Output Example

1. Performance Overview:

- Current score: 0.18666 (28% improvement from initial 0.17333)
- Inconsistent performance across iterations (scores: 0.17→0.16→0.0→0.22→0.0→0.22)
- Best performance achieved in iterations 3 and 5 with score 0.22666

2. Key Error Patterns:

A. Mathematical Reasoning Failures:

- Complex geometric problem (#70) shows systematic failure to:
 - Properly handle cyclic quadrilateral properties
 - Apply collinearity constraints
 - Use similar triangle relationships effectively
- Function intersection problem (#63) reveals:
 - Inadequate periodic function analysis
 - Poor handling of composite functions
 - Failure to systematically count intersection points

B. Framework Implementation Issues:

- Overcomplicated input processing attempts (iterations 1-2)
- Inconsistent JSON output formatting

3. Successful Elements:

- Basic mathematical structure identification
- Geometric component recognition
- Clear prompt organization with validation steps
- Strong JSON output examples

4. Improvement Pattern Analysis:

Most Effective Changes (Iterations 3/5):

- Balanced mathematical structure detection
- Clear geometric reasoning guidance
- Simple but effective input processing
- Strong output format examples

Counterproductive Changes:

- Overly complex pattern analysis (iteration 1)
- Excessive preprocessing (iteration 2)
- Oversimplified approach (iteration 4)

5. Recommended Focus Areas:

A. Input Processing:

- Keep `clean_problem` and `math_structure`
- Focus `geometric_analysis` on fundamental properties
- Avoid complex pattern detection

B. Prompt Engineering:

- Strengthen systematic proof structure
- Enhance case analysis guidance
- Maintain clear validation steps

C. Mathematical Framework:

- Add explicit proof technique selection
- Improve periodic function analysis
- Strengthen geometric relationship handling

6. Framework-Compatible Implementation:

- Use computed inputs for basic math structure [...]

I ANALYZER AND IMPROVER AGENTS: PROMPT TEMPLATES

Table 6: Analyzer and Improver agents. Temperature is set to 0.7 to encourage exploration.

Analyzer Agent
Instructions: You are the Analyzer agent in the Iterative Task Improvement Pattern that evaluates LLM agent performance on tasks by examining inputs, outputs, and ground truth. Analyze error patterns and suggest framework-compatible improvements.
LLM Agent Framework: Framework Documentation
You are an analyzer agent that works as part of this iterative improvement workflow: Iterative Improvement Documentation
Current Task Agent Configuration: Current Task Config
Current Task Agent Implementation: Current Task Implementation
Evaluation Summary: Evaluation Summary
Examples Where Task Agent Was CORRECT: Correct Examples
Examples Where Task Agent Made ERRORS: Incorrect Examples
Iteration History: Iteration History
Analysis Focus: Error patterns across multiple runs, Framework constraints (single LLM call, config-driven), General improvements (avoid overfitting to sample data), Past iteration results and performance trends, Which changes had positive/negative effects on performance, Correlation between specific improvements and score changes
Output Format: Output Format Requirements
Improver Agent
Instructions: You are the Improver agent in the Iterative Task Improvement Pattern. Generate improved task agents that follow LLM Agent Framework constraints (single LLM call, config-driven architecture). IMPORTANT: The goal is to create agents that generalize well to UNSEEN data from the same distribution. The current development data is only for evaluation - focus on improvements that will work on new, unseen problems rather than overfitting to the specific development examples.
LLM Agent Framework: Framework Documentation
Iterative Improvement Pattern: Iterative Improvement Documentation
Iteration History: Iteration History
Current Task Agent: Current Task Agent
Best Performing Agent (Score: Best Agent Score) from iteration (Best Agent Iteration): Best Task Agent
Analysis Results: Analysis Results
Improvement Guidelines: Framework Implementation Rules: All improvements must work within single LLM call architecture, Add computed inputs with processing functions for data transformation, Reference new inputs in prompt templates using {inputs.input_name}, Only implement class methods referenced by computed inputs. Explain changes like: 'Added computed input X that processes Y using function Z, referenced in prompt template as {inputs.X} to achieve improvement W'. IMPORTANT: Remember you can only modify config fields with _modifiable set to True.
Format Requirements: Output Format Requirements

J AGENT GENERATED BY SONNET IN A PRNG-BBS IMPROVEMENT LOOP

```
1458 1 {
1459 2   "description": "Enhanced Funcl solver with comprehensive arithmetic validation",
1460 3   "input_schema": "funcl_solver_input_schema.json",
1461 4   "model": {
1462 5     "_modifiable": true,
1463 6     "name": "anthropic.claude-3-5-sonnet-20241022-v2:0",
1464 7     "parameters": {
1465 8       "temperature": 0.0,
1466 9       "max_tokens": 50000
1467 10    },
1468 11  },
1469 12  "inputs": {
1470 13    "_modifiable": true,
1471 14    "function_implementation": {
1472 15      "source": "computed",
1473 16      "function": "extract_implementation",
1474 17      "args": {
1475 18        "problem_data": "raw_problem_data"
1476 19      }
1477 20    },
1478 21    "max_value_validator": {
1479 22      "source": "computed",
1480 23      "function": "max_value_validator",
1481 24      "args": {
1482 25        "problem_data": "raw_problem_data"
1483 26      }
1484 27    },
1485 28    "value_tracker": {
1486 29      "source": "computed",
1487 30      "function": "value_tracker",
1488 31      "args": {
1489 32        "problem_data": "raw_problem_data"
1490 33      }
1491 34    },
1492 35    "checksum_generator": {
1493 36      "source": "computed",
1494 37      "function": "checksum_generator",
1495 38      "args": {
1496 39        "problem_data": "raw_problem_data"
1497 40      }
1498 41    },
1499 42    "division_validator": {
1500 43      "source": "computed",
1501 44      "function": "division_validator",
1502 45      "args": {
1503 46        "problem_data": "raw_problem_data"
1504 47      }
1505 48    },
1506 49    "quotient_checker": {
1507 50      "source": "computed",
1508 51      "function": "quotient_checker",
1509 52      "args": {
1510 53        "problem_data": "raw_problem_data"
1511 54      }
1512 55    },
1513 56    "range_validator": {
1514 57      "source": "computed",
1515 58      "function": "range_validator",
1516 59      "args": {
1517 60        "problem_data": "raw_problem_data"
1518 61      }
1519 62    },
1520 63    "function_steps": {
1521 64      "source": "computed",
1522 65      "function": "extract_function_steps",
1523 66      "args": {
1524 67        "problem_data": "raw_problem_data"
1525 68      }
1526 69    },
1527 70    "hyperparams": {
1528 71      "source": "computed",
1529 72      "function": "extract_hyperparams",
1530 73      "args": {
1531 74        "problem_data": "raw_problem_data"
1532 75      }
1533 76    },
1534 77    "initial_value": {
1535 78      "source": "computed",
```

```

1512     "function": "extract_initial_value",
1513     "args": {
1514         "problem_data": "raw_problem_data"
1515     },
1516     "n": {
1517         "source": "computed",
1518         "function": "get_n_parameter",
1519         "args": {
1520             "n": 2
1521         }
1522     },
1523     "raw_problem_data": {
1524         "source": "data",
1525         "path": "$",
1526         "required": true
1527     }
1528 },
1529 "prompt": {
1530     "_modifiable": true,
1531     "sections": {
1532         "instructions": "You are a precise arithmetic calculator focused on careful validation
1533         of each calculation step. Follow ALL validation requirements exactly and show every
1534         intermediate value and check.",
1535         "task_description": "Calculate the value at position {inputs.n} using rigorous
1536         validation. Start with initial_value={inputs.initial_value}. Record and verify ALL
1537         intermediate values.",
1538         "validation_requirements": "Large Number Validation:\n{inputs.max_value_validator}\n\
1539         nValue Chain Tracking:\n{inputs.value_tracker}\n\nChecksum Validation:\n{inputs.
1540         checksum_generator}\n\nDivision Validation:\n{inputs.division_validator}\n\
1541         nQuotient Checking:\n{inputs.quotient_checker}\n\nRange Validation:\n{inputs.
1542         range_validator}",
1543         "function_steps": "Calculation Steps:\n{inputs.function_steps}",
1544         "output_format": "Respond with JSON:\n\n \n \"reasoning\": \"Show each arithmetic step
1545         with ALL intermediate values, validation checks, checksums, and cross-validations
1546         \",\n \n \"answer\": \"final numeric result as string\"\n\n"
1547     },
1548     "template": "### Instructions:\n{prompt.sections.instructions}\n\n### Task:\n{prompt.
1549     sections.task_description}\n\n### Implementation:\n{inputs.function_implementation}\n
1550     \n\n### Hyperparameters:\n{inputs.hyperparams}\n\n### Required Validation Steps:\n{
1551     prompt.sections.validation_requirements}\n\n### Calculation Process:\n{prompt.
1552     sections.function_steps}\n\n### Output Format:\n{prompt.sections.output_format}\n\
1553     nCalculate step by step, showing all validation:"
1554 },
1555 "output": {
1556     "_modifiable": false,
1557     "format": "json",
1558     "schema": {
1559         "type": "object",
1560         "required": [
1561             "reasoning",
1562             "answer"
1563         ],
1564         "properties": {
1565             "reasoning": {
1566                 "type": "string"
1567             },
1568             "answer": {
1569                 "type": "string"
1570             }
1571         }
1572     },
1573     "error_values": {
1574         "reasoning": "Error occurred during reasoning.",
1575         "answer": "PARSE_ERROR"
1576     }
1577 },
1578 "settings": {
1579     "_modifiable": true,
1580     "class_name": "Func1SolverAgent",
1581     "timeout_seconds": 300
1582 }
1583 }
1584
1585
1586
1587
1588
1589
1590
1591
1592
1593
1594
1595
1596
1597
1598
1599
1600
1601
1602
1603
1604
1605
1606
1607
1608
1609
1610
1611
1612
1613
1614
1615
1616
1617
1618
1619
1620
1621
1622
1623
1624
1625
1626
1627
1628
1629
1630
1631
1632
1633
1634
1635
1636
1637
1638
1639
1640
1641
1642
1643
1644
1645
1646
1647
1648
1649
1650
1651
1652
1653
1654
1655
1656
1657
1658
1659
1660
1661
1662
1663
1664
1665
1666
1667
1668
1669
1670
1671
1672
1673
1674
1675
1676
1677
1678
1679
1680
1681
1682
1683
1684
1685
1686
1687
1688
1689
1690
1691
1692
1693
1694
1695
1696
1697
1698
1699
1700
1701
1702
1703
1704
1705
1706
1707
1708
1709
1710
1711
1712
1713
1714
1715
1716
1717
1718
1719
1720
1721
1722
1723
1724
1725
1726
1727
1728
1729
1730
1731
1732
1733
1734
1735
1736
1737
1738
1739
1740
1741
1742
1743
1744
1745
1746
1747
1748
1749
1750
1751
1752
1753
1754
1755
1756
1757
1758
1759
1760
1761
1762
1763
1764
1765
1766
1767
1768
1769
1770
1771
1772
1773
1774
1775
1776
1777
1778
1779
1780
1781
1782
1783
1784
1785
1786
1787
1788
1789
1790
1791
1792
1793
1794
1795
1796
1797
1798
1799
1800
1801
1802
1803
1804
1805
1806
1807
1808
1809
1810
1811
1812
1813
1814
1815
1816
1817
1818
1819
1820
1821
1822
1823
1824
1825
1826
1827
1828
1829
1830
1831
1832
1833
1834
1835
1836
1837
1838
1839
1840
1841
1842
1843
1844
1845
1846
1847
1848
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
1920
1921
1922
1923
1924
1925
1926
1927
1928
1929
1930
1931
1932
1933
1934
1935
1936
1937
1938
1939
1940
1941
1942
1943
1944
1945
1946
1947
1948
1949
1950
1951
1952
1953
1954
1955
1956
1957
1958
1959
1960
1961
1962
1963
1964
1965
1966
1967
1968
1969
1970
1971
1972
1973
1974
1975
1976
1977
1978
1979
1980
1981
1982
1983
1984
1985
1986
1987
1988
1989
1990
1991
1992
1993
1994
1995
1996
1997
1998
1999
2000

```

```

1566 7 sys.path.append(os.path.join(os.path.dirname(__file__), '..', '..', '..', '..', 'llm_agent'))
1567 8
1568 9 from llm_agent import LLMAgent
1569 10
1570 11 class Func1SolverAgent(LLMAgent):
1571 12     """Enhanced Func1 solver with arithmetic validation."""
1572 13
1573 14     def extract_implementation(self, problem_data):
1574 15         """Extract function implementation from problem data."""
1575 16         return problem_data.get('implementation', '')
1576 17
1577 18     def extract_hyperparams(self, problem_data):
1578 19         """Extract and format hyperparameters."""
1579 20         hyperparams = problem_data.get('hyperparams', {})
1580 21         if isinstance(hyperparams, dict):
1581 22             formatted = []
1582 23             for key, value in hyperparams.items():
1583 24                 formatted.append(f"{key}: {value}")
1584 25             return "\n".join(formatted)
1585 26         return str(hyperparams)
1586 27
1587 28     def extract_initial_value(self, problem_data):
1588 29         """Extract initial value from problem data."""
1589 30         return str(problem_data.get('initial_value', ''))
1590 31
1591 32     def get_n_parameter(self, n):
1592 33         """Return the position parameter."""
1593 34         return str(n)
1594 35
1595 36     def max_value_validator(self, problem_data):
1596 37         """Generate validation steps for large number arithmetic."""
1597 38         steps = [
1598 39             "Large Number Validation:",
1599 40             "1. For multiplication A * B:",
1600 41             "    a. Split into single digits: A = a a ... a , B = b b ... b ",
1601 42             "    b. Calculate partial products:",
1602 43             "        - For each digit pair ( a , b ): p = a * b ",
1603 44             "        - Track carries explicitly",
1604 45             "    c. Sum partial products with position tracking",
1605 46             "    d. Verify: result = sum of all partial products",
1606 47             "2. Cross-validate using reverse calculation:",
1607 48             "    - Split result into chunks",
1608 49             "    - Verify chunks match partial products"
1609 50         ]
1610 51         return "\n".join(steps)
1611 52
1612 53     def value_tracker(self, problem_data):
1613 54         """Generate value tracking steps for calculation chain integrity."""
1614 55         steps = [
1615 56             "Value Chain Tracking:",
1616 57             "1. For each position i in calculation:",
1617 58             "    a. Record value[i] = current value",
1618 59             "    b. Record squared[i] = value[i] ** 2",
1619 60             "    c. Record quotient[i] = squared[i] // modulus",
1620 61             "    d. Record remainder[i] = value[i+1]",
1621 62             "2. Verify chain integrity:",
1622 63             "    - remainder[i] = squared[i] - (quotient[i] * modulus)",
1623 64             "    - 0 < remainder[i] < modulus",
1624 65             "    - value[i+1] = remainder[i]"
1625 66         ]
1626 67         return "\n".join(steps)
1627 68
1628 69     def checksum_generator(self, problem_data):
1629 70         """Generate checksum validation steps between calculations."""
1630 71         steps = [
1631 72             "Checksum Validation:",
1632 73             "1. For each step i:",
1633 74             "    a. Calculate checksum[i] = (value[i] * 31 + squared[i]) mod 997",
1634 75             "    b. Verify: checksum[i] matches independent calculation",
1635 76             "2. Cross-step validation:",
1636 77             "    - Verify: checksum[i+1] = (remainder[i] * 31) mod 997",
1637 78             "    - Track checksum chain for each position"
1638 79         ]
1639 80         return "\n".join(steps)
1640 81
1641 82     def division_validator(self, problem_data):
1642 83         """Generate enhanced division validation steps."""
1643 84         modulus = problem_data.get('hyperparams', {}).get('modulus', 0)
1644 85         steps = [
1645 86             f"Division Validation Steps:",
1646 87             f"1. For dividing number A by {modulus}:",

```

```

1620         f"    a. Calculate quotient  $Q = A \mod \{modulus\}$ ",
1621         f"    b. Verify:  $Q \mod \{modulus\} = A$ ",
1622         f"    c. Verify:  $(Q + 1) \mod \{modulus\} > A$ ",
1623         f"    d. Record  $Q$  for remainder calculation",
1624         f"    e. Track sub-steps:",
1625         f"        - Record partial quotients",
1626         f"        - Verify digit-by-digit multiplication",
1627         f"        - Cross-validate remainder"
1628     ]
1629     return "\n".join(steps)
1630
1631 def quotient_checker(self, problem_data):
1632     """Generate enhanced quotient checking steps."""
1633     steps = [
1634         "Quotient Verification:",
1635         "1. Record squared value (S)",
1636         "2. Calculate and verify quotient (Q):",
1637         "    -  $Q = S \mod \{modulus\}$  (integer division)",
1638         "    - Verify:  $Q \mod \{modulus\} = S$ ",
1639         "    - Verify:  $(Q + 1) \mod \{modulus\} > S$ ",
1640         "    - Perform digit-by-digit multiplication check",
1641         "3. Calculate remainder:",
1642         "    -  $R = S - (Q \mod \{modulus\})$ ",
1643         "4. Verify remainder is correct:",
1644         "    -  $0 \leq R < \{modulus\}$ ",
1645         "    - Cross-validate with value chain"
1646     ]
1647     return "\n".join(steps)
1648
1649 def range_validator(self, problem_data):
1650     """Generate enhanced range validation steps."""
1651     modulus = problem_data.get('hyperparams', {}).get('modulus', 0)
1652     steps = [
1653         f"Range Validation:",
1654         f"1. For each position i:",
1655         f"    - Verify value[i]  $\geq 0$ ",
1656         f"    - Verify value[i] < {modulus}",
1657         f"2. For squared values at position i:",
1658         f"    - Record exact value squared[i]",
1659         f"    - Verify squared[i] = value[i] * value[i]",
1660         f"    - Verify using digit-by-digit multiplication",
1661         f"3. Position-specific bounds:",
1662         f"    - Track maximum possible value at each position",
1663         f"    - Verify values stay within position bounds"
1664     ]
1665     return "\n".join(steps)
1666
1667 def extract_function_steps(self, problem_data):
1668     """Extract and format key arithmetic operations."""
1669     steps = [
1670         "1. Initialize calculation chain:",
1671         "    - Set value[0] = initial_value",
1672         "    - Calculate checksum[0]",
1673         "2. For each position i:",
1674         "    a. Square Operation:",
1675         "        - Calculate value[i] using digit-by-digit multiplication",
1676         "        - Record squared[i] and partial products",
1677         "        - Verify checksum[i]",
1678         "    b. Division and Modulo:",
1679         "        - Calculate quotient[i] with validation",
1680         "        - Calculate remainder[i] = value[i+1]",
1681         "        - Verify chain integrity",
1682         "    c. Cross-validate:",
1683         "        - Verify all checksums match",
1684         "        - Confirm value chain consistency",
1685         "        - Validate position-specific bounds"
1686     ]
1687     return "\n".join(steps)

```

K ACCURACY-DRIVEN IMPROVEMENT LOOPS

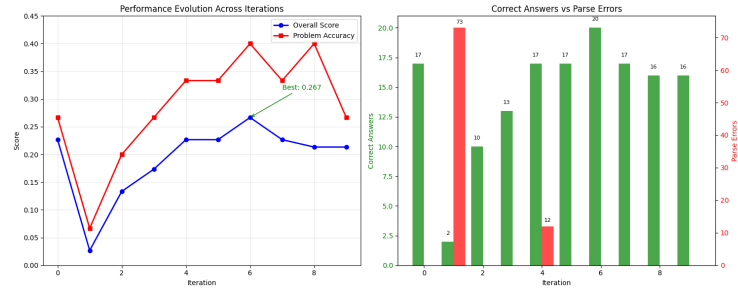


Figure 2: Aime 2024 improvement loop using Sonnet 3.5 v2 as meta agents.

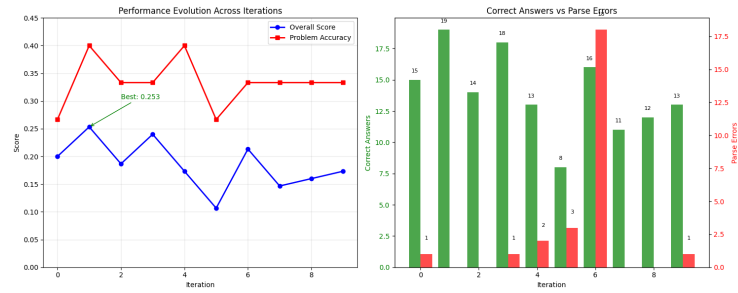


Figure 3: Aime 2024 improvement loop using OpenAI 120b reasoning model as meta agents.

L AGENT GENERATED BY ACCURACY+COST OPTIMIZATION LOOP BY GPT-OSS-120B

```
1728 1 "description": "AIME 2024 ICL task agent \u2013 focused rhombus example, explicit bound,
1729 2   streamlined prompt",
1730 3   "input_schema": "aime2024_icl_solver_input_schema.json",
1731 4   "model": {
1732 5     "_modifiable": true,
1733 6     "name": "anthropic.claude-3-5-sonnet-20241022-v2:0",
1734 7     "parameters": {
1735 8       "temperature": 0.2,
1736 9       "max_tokens": 500
1737 10    },
1738 11  },
1739 12  "inputs": {
1740 13    "_modifiable": true,
1741 14    "problem_text": {
1742 15      "source": "computed",
1743 16      "function": "extract_problem_text",
1744 17      "args": {
1745 18        "problem_data": "raw_problem_data"
1746 19      }
1747 20    },
1748 21    "rhombus_property": {
1749 22      "source": "computed",
1750 23      "function": "rhombus_property",
1751 24      "args": {}
1752 25    },
1753 26    "solve_perpendicular_condition": {
1754 27      "source": "computed",
1755 28      "function": "solve_perpendicular_condition",
1756 29      "args": {
1757 30        "problem_data": "raw_problem_data"
1758 31      }
1759 32    },
1760 33    "bd2_bound": {
1761 34      "source": "computed",
1762 35      "function": "bd2_bound",
1763 36      "args": {}
1764 37    },
1765 38    "rhombus_icl_example": {
1766 39      "source": "computed",
1767 40      "function": "rhombus_icl_example",
1768 41      "args": {}
1769 42    },
1770 43    "guard_text": {
1771 44      "source": "computed",
1772 45      "function": "guard_text",
1773 46      "args": {}
1774 47    },
1775 48    "raw_problem_data": {
1776 49      "source": "data",
1777 50      "path": "$",
1778 51      "required": true
1779 52    }
1780 53  },
1781 54  "prompt": {
1782 55    "_modifiable": true,
1783 56    "sections": {
1784 57      "instructions": "You are a mathematical problem solver. Solve the given AIME problem and
1785 58      output the answer as a three\u2011digit integer.",
1786 59      "output_format": "You must respond with a JSON object containing exactly two fields:\n-
1787 60      \"reasoning\": string with your detailed solution\n- \"answer\": string with a
1788 61      three\u2011digit answer (pad with leading zeros if needed)\n\nFormat: {\"reasoning
1789 62      \": \"Your detailed solution here\", \"answer\": \"XXX\"}\n\nExample:\n{\"reasoning
1790 63      \": \"To solve..., I ...\", \"answer\": \"042\"}"
1791 64    },
1792 65    "template": "### Instructions:\n{prompt.sections.instructions}\n\n### Geometry hint:\n{
1793 66      inputs.rhombus_property}\n\n### Algebraic bound hint:\n{inputs.
1794 67      solve_perpendicular_condition}\n\n### Known bound value:\nThe minimal possible BD\
1795 68      u00b2 is {inputs.bd2_bound}.\n\n### Example to follow:\n{inputs.rhombus_icl_example}\
1796 69      \n\n### Problem:\n{inputs.problem_text}\n\n### Output format:\n{prompt.sections.
1797 70      output_format}\n\n{inputs.guard_text}"
1798 71  },
1799 72  "output": {
1800 73    "_modifiable": false,
1801 74    "format": "json",
1802 75    "schema": {
1803 76      "type": "object",
```

```

1782 66     "required": [
1783 67         "reasoning",
1784 68         "answer"
1785 69     ],
1786 70     "properties": {
1787 71         "reasoning": {
1788 72             "type": "string"
1789 73         },
1790 74         "answer": {
1791 75             "type": "string",
1792 76             "pattern": "[0-9]{3}$"
1793 77         }
1794 78     },
1795 79     "error_values": {
1796 80         "reasoning": "Error occurred during reasoning.",
1797 81         "answer": "PARSE_ERROR"
1798 82     }
1799 83 },
1800 84     "settings": {
1801 85         "_modifiable": false,
1802 86         "class_name": "AIME2024ICLSolverAgent",
1803 87         "timeout_seconds": 300,
1804 88         "exit_on_parse_failure": false
1805 89     }
1806 90 }
1807 91 }

```

```

1801 1  from llm_agent import LLMAgent
1802 2  from typing import Dict, Any
1803 3
1804 4  class AIME2024TaskAgent(LLMAgent):
1805 5      """
1806 6      AIME2024-specific reasoning + answer agent.
1807 7
1808 8      Processes AIME mathematical competition problems and generates detailed step-by-step
1809 9      reasoning followed by answers.
1810 10     """
1811 11     def __init__(self, name="AIME2024_Task", config_path=None, dry_run=False):
1812 12         super().__init__(name, config_path, dry_run)
1813 13
1814 14     def extract_problem_text(self, problem_data: Dict[str, Any]) -> str:
1815 15         """
1816 16         Extract problem text from AIME2024 data format.
1817 17
1818 18         AIME2024 format:
1819 19         {
1820 20             "id": int,
1821 21             "problem": str, # Main problem statement
1822 22             "url": str
1823 23         }
1824 24         """
1825 25         return str(problem_data.get('problem', ''))

```

1819
1820
1821
1822
1823
1824
1825
1826
1827
1828
1829
1830
1831
1832
1833
1834
1835