

Experience-Guided Adaptation of Inference-Time Reasoning Strategies

Adam Stein^{1*}, Matthew Trager², Benjamin Bowman², Michael Kleinman²,
Aditya Chattopadhyay², Wei Xia², Stefano Soatto²

¹University of Pennsylvania, ²AWS AI

Abstract

Enabling agentic AI systems to adapt their problem-solving approaches based on post-training interactions remains a fundamental challenge. While systems that update and maintain a memory at inference time have been proposed, existing designs only steer the system by modifying textual input to a language model or agent, which means that they cannot change sampling parameters, remove tools, modify system prompts, or switch between agentic and workflow paradigms. On the other hand, systems that adapt more flexibly require offline optimization and remain static once deployed. We present Experience-Guided Reasoner (EGUR), which generates tailored *strategies*—complete computational procedures involving LLM calls, tools, sampling parameters, and control logic—dynamically at inference time based on accumulated experience. We achieve this using an LLM-based *meta-strategy*—a strategy that outputs strategies—enabling adaptation of all strategy components (prompts, sampling parameters, tool configurations, and control logic). EGUR operates through two components: a **Guide** generates multiple candidate strategies conditioned on the current problem and structured memory of past experiences, while a **Consolidator** integrates execution feedback to improve future strategy generation. This produces complete, ready-to-run strategies optimized for each problem, which can be cached, retrieved, and executed as needed without wasting resources. Across five challenging benchmarks (AIME 2025, 3-SAT, and three Big Bench Extra Hard tasks), EGUR achieves up to 14% accuracy improvements over the strongest baselines while reducing computational costs by up to 111×, with both metrics improving as the system gains experience.

1 Introduction

Modern AI systems employ sophisticated *strategies*—complex computational procedures involving LLM calls, tools, and control logic, often implemented using frameworks like SGLang [1] or DSPy [2]—to tackle challenging reasoning tasks. For example, systems solving International Mathematical Olympiad (IMO) problems dynamically decompose problems into subgoals, verify solutions through parallel sampling, and iteratively self-correct [3, 4]. Similarly, code generation agents switch between different prompting approaches, adjust sampling temperatures, and selectively use tools based on problem complexity [2, 5].

While these strategies are often designed to be general-purpose and highly adaptive to inputs—particularly agents, which consist of loops of LLM and tool execution—they generally remain *static* at inference time, unable to learn from experience. For example, without persistent memory, an AI system for mathematical reasoning will apply the same expensive multi-step decomposition every time it encounters a problem it has solved in the past (potentially even failing to solve it again), and will repeatedly make identical mistakes on problems it has failed to solve before.

To address this problem, different methods for adapting AI systems have been proposed. Approaches such as Dynamic Cheatsheet [6] and Buffer of Thoughts [7] maintain memory across problems, but this memory is used only for *textual steering* of an LLM or agent. As such, they cannot change sampling parameters, add or remove tools, modify control logic, or cache successful computational procedures for direct

*Work done during an internship at AWS AI. Correspondence to: Matthew Trager <mttrager@amazon.com>.

reuse. On the other hand, offline methods like ADAS [8] can adapt strategies much more flexibly, but require expensive training phases and remain static once deployed, unable to continue learning from new experiences. The core challenge is to **adapt the computational structure of strategies from experience, to continually improve both accuracy and efficiency.**

To address this challenge, we present Experience-Guided Reasoner (EGUR), a system designed to generate strategies dynamically at inference time. Our approach uses an LLM-based *meta-strategy*—a strategy that outputs strategies—rather than steering existing strategies at runtime, it proposes new strategies. This enables generation of complete, tailored computational procedures for each problem that can be easily conditioned using a memory store and enables simple and effective caching of successful strategies for reuse. EGUR operates through two main components: a **Guide** that generates multiple candidate strategies conditioned on the current problem and structured memory of past experiences, and a **Consolidator** that processes execution outcomes to improve future strategy generation. By generating and comparing multiple strategies per problem, the system learns the relative effectiveness of different strategies, leading to continual improvements in both accuracy and computational efficiency.

Our contributions are:

- We introduce a system that generates strategies dynamically at inference time based on accumulated experience, enabling adaptation of all strategy components including prompts, sampling parameters, tools, and control logic (Section 3).
- We formalize strategies as compositions of stateful *processes* which map inputs to outputs while also updating a state. This provides a unified representation for diverse reasoning paradigms (workflows, agents, tool-use systems) that supports compositional cost tracking and execution tracing (Section 2.1).
- We demonstrate that generating multiple strategies per problem and comparing their effectiveness enables continual improvement in both accuracy and efficiency (Figure 5).
- Across five challenging benchmarks (AIME 2025, 3-SAT, and three Big Bench Extra Hard tasks), EGUR achieves up to 14% accuracy improvements (EGUR vs. Mem0 on 3-SAT) while reducing computational costs by up to 111x (EGUR vs. DC on Object Counting), with both metrics improving as the system gains experience (Figure 4).

2 Strategies For AI Systems

In this section, we formalize strategies as compositions of stateful processes—functions taking inputs and a state and producing an output and an updated state—providing a unified way to describe any strategy in terms of its components. This yields a natural taxonomy of common strategies and enables us to describe feedback mechanisms as state-modification methods.

2.1 Representing Strategies

We use the term *strategy* for a specification of the computational procedure that an AI system applies to inputs to produce outputs, capturing LLM inference calls, tool invocations, and control flow decisions. For example, the Chain-of-Thought (CoT) strategy involves a single LLM call with specific prompting, while CodeAct iteratively calls an LLM and executes generated code until a termination condition is met.

An adaptive strategy engages in a sequence of interactions, or episodes, and updates its internal state based on past outcomes. To enable adaptive strategies, we need a representation that explicitly captures both what can be adapted and how state evolves across episodes. Existing frameworks like DSPy [2] and SGLang [1] implement strategies as Python code without explicitly marking adaptable components. TextGrad [5] and Trace [10] mark components for optimization, but are designed for offline training and lack explicit state management for continual learning during inference. We introduce a compositional formalism that unifies adaptable components with state management, enabling meta-strategies to generate and refine strategies from accumulated experience.

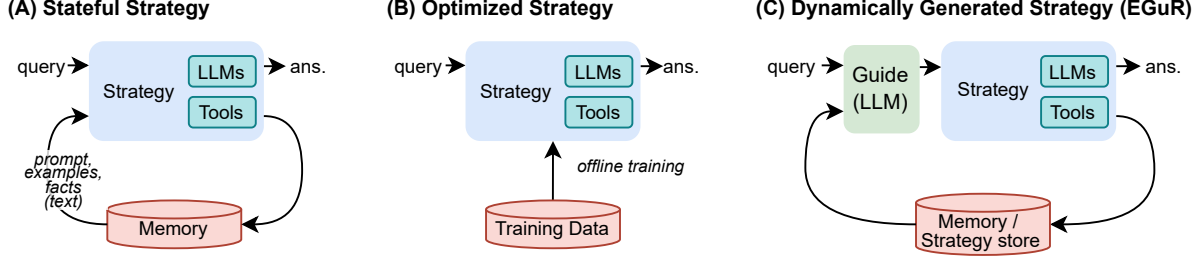


Figure 1: Comparison of existing experience-based adaptation methods (A and B) to our approach in (C) which dynamically produces a compiled strategy based on the current query and memory at inference time. Methods which augment existing strategies with state, such as Dynamic Cheatsheet[6] and Mem0[9], are depicted in (A) and methods such as ADAS [8] and OPTO [10] which optimize strategies offline are shown in (B). Our method shown in (C) uses a state (which can contain useful compiled strategies) during inference to guide the system in producing effective strategies for each query, unlike existing methods which cannot adapt their strategies per-query.

We model strategies as compositions of stateful *processes*. A process is a function that takes an input and a state (e.g., conversation history, sampling parameters), produces an output, and updates the state. This explicit state management clarifies what components can be adapted (prompts, temperature, available tools). Formally, we write a process type as $A \xRightarrow{\Sigma} B : (A \times \Sigma) \rightarrow (B \times \Sigma)$ where A and B are input and output domains and Σ is the state space. Strategies are built compositionally using the grammar shown in Figure 2. Sequential composition ($S_1; S_2$) chains processes, parallel composition ($S_1 \parallel S_2$) executes concurrently, conditionals enable branching, and recursion enables iteration.

For instance, CodeAct (shown in Figure 2b) composes an LLM call with a code execution tool recursively, iterating until the LLM output contains no code. The LLM call process maintains conversation history in its state, while the code execution process maintains the execution environment.

This compositional structure is key to our approach. By explicitly representing all components (prompts, sampling parameters, tools, control flow) and their state dependencies, we enable a meta-strategy to generate new strategies by composing different components. For instance, our meta-strategy can adapt CodeAct by changing the LLM’s prompt, adjusting temperature, or removing the code interpreter entirely for non-algorithmic problems. This framework also allows us to formalize important properties of strategies: computational cost and intermediate traces (Appendix B.4), formal semantics (Appendix B.2), and structural characteristics (discussed in the next section).

2.2 The Strategy Landscape

We organize strategies into several categories based on their structure. A strategy is *dynamic* if it contains conditionals, meaning the control flow can change at runtime based on the input. A dynamic strategy that is recursive (uses **recfun** as shown in Figure 2a) is called an *agent*. Strategies that do not use recursion are called *workflows* since their runtime behavior is less dynamic, and a workflow that is not dynamic is called a *pipeline*. Strategies can also use parallelization to solve several tasks concurrently or to produce multiple responses at once, which can improve accuracy and reliability [11]. Table 1 shows a breakdown of five commonly used strategies based on their use of parallel compute, their dynamic or agentic nature, and their tool utilization. A description of the five strategies is in Appendix B.6. Our classification differs from common definitions that do not classify Eval-Opt as agentic and lack a clear way to determine when a strategy is agentic [12].

Figure 3 shows that different strategies result in significantly different behavior in terms of cost and accuracy. Among strategies we consider, the Code strategy is best for 3-SAT and Word Sorting, but worst for AIME and Movie Recommendation. Even among strategies with comparable accuracy, cost can vary dramatically: Eval-Opt generally achieves similar accuracy to Self-Consistency but often at half the cost. Notably, more expressive strategies do not guarantee better performance: while agentic strategies can the-

$S ::= \mathbf{base}_p$	base process	
$S_1; S_2$	sequential composition where	
	$S_1 : A \xRightarrow{\Sigma_1} C$ and $S_2 : C \xRightarrow{\Sigma_2} B$	
	and Σ_1, Σ_2 are views of Σ	
$S_1 \parallel S_2$	parallel composition where	
	$S_1 : A \xRightarrow{\Sigma_1} B_1$ and $S_2 : A \xRightarrow{\Sigma_2} B_2$ and	
	$B = (B_1, B_2)$ and Σ_1, Σ_2 are disjoint views of Σ	
$\mathbf{if } S_1 \mathbf{ then } S_2 \mathbf{ else } S_3$	conditional where $S_1 : A \xRightarrow{\Sigma_1} \text{Bool}$,	
	$S_2 : A \xRightarrow{\Sigma_2} B$, and $S_3 : A \xRightarrow{\Sigma_3} B$	
	and $\Sigma_1, \Sigma_2, \Sigma_3$ are views of Σ	
$\mathbf{recfun } f : S_1$	recursive definition of process f where	
	$S_1 : A \xRightarrow{\Sigma_1} B$ which can internally use f	
	and Σ_1 is a view of Σ .	

(a) Grammar for strategies where $S : A \xRightarrow{\Sigma} B$.

```

1 recfun CodeAct :
2   CallLLM;
3   if ContainsCode
4   then (ExecCode; CodeAct)
5   else return

```

(b) Example: CodeAct Strategy

Figure 2: High-level grammar for strategies as compositions of base processes (left) with an example of how CodeAct is represented in this grammar (right). A further formalization is provided in Appendix B where we define how to additionally construct processes from pure functions as well as processes for explicitly accessing and updating the state, and some useful base processes such as **return**.

oretically emulate simpler strategies through their conditional and recursive structure, they may fail to select the right behavior in practice and, even when successful, often incur substantially higher computational costs than using the simpler strategy directly (for example, CodeAct is not the most effective strategy in Figure 3, despite being the most general). These observations highlight a fundamental trade-off where practitioners must balance solution quality against computational expense, often termed “overthinking.”

2.3 Incorporating Feedback to Strategies

The observations above motivate systems that can adapt their strategy selection to each problem. Ideally, such systems should improve with experience: using information from past problem-solving attempts to make better strategy choices on future problems. This requires maintaining a state during inference that captures which strategies work well for which types of problems.

Existing work addresses two separate problems but not both simultaneously. Some work focuses on maintaining state during testing but is limited in what the state can influence. These methods add state to strategies compiled before inference, so they are limited to steering strategies by changing their input [6, 9]. Other work adapts general strategies to experience but requires expensive offline optimization to compile an optimized strategy that is then left stateless during inference [8, 10]. Figure 1 illustrates how training-based optimization methods (b) compare to state-based adaptation approaches (a) and our approach (c).

The success of offline strategy optimization in modifying entire strategies motivates a method that compiles useful strategies on the fly by learning from experience. This is the approach we take and introduce in the following section.

3 Experience-Guided Reasoner (EGUR)

In this section, we introduce a system that generates tailored strategies dynamically at inference time by learning from accumulated experience. Our approach centers on two key components: a *Guide* that gen-

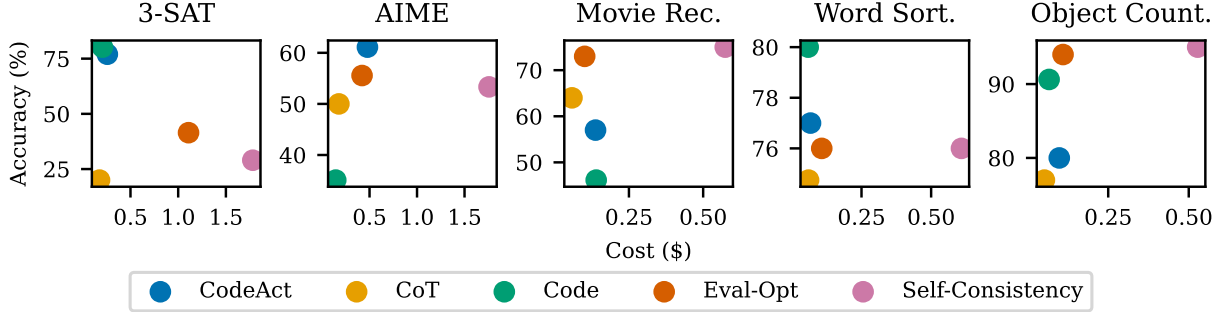


Figure 3: Comparison of strategy performance across tasks for Claude 3.7 Sonnet. Strategies closer to the top-left corner are best for the task in terms of accuracy and cost. The optimal strategy differs significantly across tasks. For example, Code excels on 3-SAT and Word Sorting but performs poorly on Movie Recommendation and AIME. Descriptions of the above strategies are provided in Appendix B.6.

Table 1: Commonly used strategies in terms of their used design patterns and tools. CI stands for code interpreter. Pipelines do not use conditionals, workflows are any strategy using conditionals, and agents are strategies using recursion.

Name	Parallelization	Conditionals	Recursion	Tools (base processes)
Pipelines	–	✗	–	–
CoT [13]				LLM
Self-Consistency [11]	✓			LLM
Workflows	–	✓	–	–
Code [14]		✓		LLM, CI
Agents	–	–	✓	–
Eval-Opt [15]		✓	✓	LLM
CodeAct [16]		✓	✓	LLM, CI

erates complete, problem-specific strategies based on past experience, and a *Consolidator* that processes execution outcomes to improve future strategy generation.

3.1 General Design

The **Guide** generates complete strategy specifications tailored to each problem based on the current context and past experiences. Rather than making decisions step-by-step like traditional agents, it produces a complete computational procedure upfront—specifying LLM calls, sampling parameters, tools, and control flow. The **Consolidator** processes execution outcomes, including reasoning traces and verifier feedback, maintaining a structured memory that improves strategy generation.

Formally, these components have types:

$$\text{Guide} : \text{Str} \xrightarrow{\Sigma} \left(\text{Str} \xrightarrow{\sigma} \text{Str} \right) \quad \text{Consolidator} : \text{List}[\text{Exp}] \xrightarrow{\Sigma} \Sigma$$

where Exp represents (question, answer, trace, cost, feedback) tuples. The state space is divided into Σ (between-episode memory maintained by Consolidator) and σ (within-episode memory used during strategy execution). The Guide has a state Σ , takes a string as input and outputs a strategy as defined in Section 2.1 (which itself uses a state σ and takes a string as input and outputs a string). The Consolidator updates Σ based on execution experiences.

The complete EGUR process is shown in Algorithm 1, and we also provide an equivalent definition using our strategy language in Appendix B.5.

Algorithm 1 Experience-Guided Reasoner

Require: Question q , context Σ , and exploration factor k .

Ensure: Answer a and updated context Σ' .

```
1:  $\{S_1, \dots, S_k\} \leftarrow \mathbf{Guide}(q, \Sigma, k)$  // generate  $k$  candidate strategies
2:  $e = []$  // initialize experience collection
3: for  $i = 1$  to  $k$  in parallel do
4:    $\sigma \leftarrow \text{null}$  // initialize strategy state
5:    $(a_i, t_i, c_i) \leftarrow \mathbf{Execute}(S_i, q, \sigma)$  // execute strategy  $S_i$ 
6:    $f_i \leftarrow \text{verifier}(q, a_i)$  // get verifier feedback
7:    $e \leftarrow e + [(q, a_i, t_i, c_i, f_i)]$  // collect experience tuple
8: end for
9:  $\Sigma' \leftarrow \mathbf{Consolidator}(e, \Sigma)$  // update context with experiences
10:  $a \leftarrow a_1$  // use answer from first strategy
11: return  $a, \Sigma'$ 
```

3.2 Strategy Selection with the Guide

The Guide generates k candidate strategies for each problem by conditioning on the query q and relevant experiences from the context Σ . This design provides some key benefits: **observability** of the system’s behavior with complete strategy specifications; **exploration** of the space of possible strategies by generating multiple candidates; and **comparison** of the relative performance of different approaches for future strategy generation. The exploration factor k balances computational cost with strategy diversity—setting $k = 1$ minimizes overhead, while larger values enable discovery of better strategies at increased cost.

We invoke the Guide k times on problem q to obtain k candidate strategies $\{S_1, \dots, S_k\}$, which we execute to obtain performance measures. Similar to Group Relative Policy Optimization (GRPO) [17], we use relative comparisons within each group of k strategies: the Consolidator records which strategies perform better relative to alternatives on similar problems, enabling the Guide to improve over time. In practice, the Guide is implemented as an LLM call that produces strategy code, which is then parsed from the response and executed (see Appendix F.1 for the prompt).

3.3 Feedback Integration with the Consolidator

The Consolidator maintains structured memory to guide future strategy generation. Rather than accumulating raw experiences, it performs selective abstraction to maximize information utility while preventing memory bloat. The Consolidator processes experiences $e = [(q_1, a_1, t_1, c_1, f_1), \dots, (q_k, a_k, t_k, c_k, f_k)]$ from executed strategies, where each tuple contains the query, answer, execution trace, computational cost, and verifier feedback.

The context Σ maintains two key components: a **Strategy Library** that stores successful strategies with their problem characteristics for reuse and template generation, and **General Notes** that capture high-level insights about strategy effectiveness, common failure patterns, useful techniques, and problem-solving heuristics. The entire context is passed directly to the Guide for each problem, enabling it to perform retrieval and synthesis in-context rather than requiring external retrieval mechanisms.

To prevent unbounded memory growth, the Consolidator implements selective retention policies, prioritizing recent experiences and frequently accessed patterns while removing outdated information to maintain responsiveness to evolving problem distributions while preserving valuable learned knowledge. In practice, the Consolidator is implemented as an LLM call that produces memory insertions and deletions which are parsed and applied (see Appendix F.2 for the prompt).

4 Experiments

We evaluate EGUR across diverse reasoning tasks to answer four key research questions: (RQ1) Does it outperform baselines in accuracy and cost? (RQ2) Does dynamic strategy generation improve over existing

stateful methods? (RQ3) Is comparative strategy evaluation important for continual improvements? (RQ4) Does EGUR learn novel and useful strategies from experience?

4.1 Experimental Setup

Datasets. We evaluate on five datasets: AIME (2022-2024 for training, 2025 held-out), 3-SAT (5-40 variables, 400 training/40 test samples), and three BBEH [18] tasks (movie recommendation, word sorting, object counting, 100 training/100 test each). We process training data in shuffled batches of ten samples, evaluating all ten in parallel then sequentially updating based on each sample’s experience.

Models. Claude 3.7 Sonnet, Qwen3-Next-80B-A3B-Thinking, and GPT-OSS-120B.

Baselines. CodeAct [16] (stateless code-interpreter agent), Dynamic Cheatsheet [6] (CodeAct with string-based memory appended to inputs), and Mem0 (CodeAct with vector database memory). We modify Dynamic Cheatsheet to incorporate verifier feedback during updates.

Verifiers. For all tasks except for 3-SAT, we use ground-truth based verifiers that compare the final answer against the correct solution. For 3-SAT, we verify the proposed assignment satisfies all clauses using a satisfiability checker. The verifier provides only binary feedback (correct/incorrect) which is passed to the Consolidator.

Implementation details. All LLM calls use temperature 0 and a maximum token budget of 20,000 tokens. For Claude 3.7 Sonnet, we enable extended thinking mode with a maximum thinking budget of 10,000 tokens; when thinking mode is enabled, Claude’s temperature is automatically set to 1.0. These settings apply to all strategy executions, Guide calls, and Consolidator calls across all methods. With EGUR, the strategy can include LLM calls without thinking mode or with non-zero temperature.

Evaluation. We report prequential accuracy (accuracy before updating on each sample) and cumulative cost for strategy execution (excluding system feedback and update costs). Continual updates use one seed; held-out evaluation uses three seeds. Costs are in USD based on API pricing.

4.2 RQ1: Does EGUR outperform baselines in accuracy and cost?

Figure 4 shows EGUR’s performance on held-out evaluation sets as training progresses. Two key trends emerge: accuracy increases with more experience, and cost decreases as EGUR learns more efficient strategies. Notably, Dynamic Cheatsheet’s costs are significantly higher than other approaches (often exceeding \$1.00 per sample) due to the increasing context length from accumulating memory, and the underlying strategy taking many turns before outputting a final answer. In contrast, EGUR maintains low and decreasing costs through selective memory management. Mem0 shows minimal cost reduction and modest accuracy improvements, suggesting limited adaptation capability.

Table 3 provides detailed prequential accuracy and cost results across all datasets and models. EGUR consistently achieves the best accuracy-cost trade-offs: on Claude 3.7 Sonnet, EGUR-5 achieves 96.0% accuracy on 3-SAT at \$0.152 cost versus CodeAct’s 77.0% at \$0.257 and Dynamic Cheatsheet’s 89.9% at \$76.353. Similar patterns hold across models and tasks.

4.3 RQ2: Does dynamic strategy generation improve over existing stateful strategies?

The key distinction between EGUR and existing stateful methods is architectural: **EGUR generates a complete strategy specification for each problem before execution**, while existing methods **steer dynamic strategies at runtime by modifying their inputs**. Concretely, Mem0 and Dynamic Cheatsheet prepend memory to a fixed CodeAct agent’s inputs, influencing its behavior indirectly. In contrast, EGUR produces entirely new strategy specifications—including prompts, sampling parameters (temperature, max tokens), tool availability, and control flow—tailored to each problem.

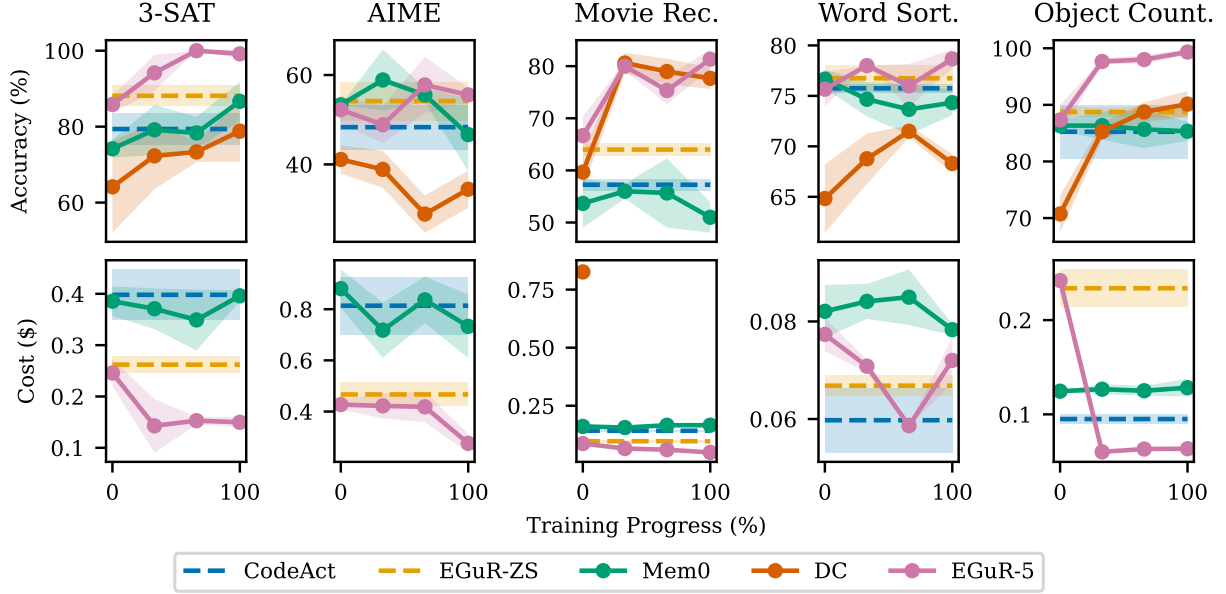


Figure 4: Evolution of accuracy and cost on held-out evaluation sets as training progresses for Claude 3.7 Sonnet. EGuR-5 consistently improves accuracy while reducing cost with more experience. Cost is shown up to \$1.0 for visualization; Dynamic Cheatsheet (DC) typically exceeds this threshold, reaching \$9.95, \$2.26, \$2.88, \$4.32, and \$7.16 per sample after training on 3-SAT, AIME, Movie Rec., Word Sort., and Object Count., respectively. The complete results for Claude 3.7 Sonnet, GPT-OSS-120B, and Qwen3-Next-80B-A3B-Thinking are included in Appendix C.

This architectural difference has profound implications. Methods that steer via input modification cannot remove tools, change sampling parameters, or switch between agentic and workflow paradigms. EGuR is significantly more flexible: for instance, it learns to completely remove the code interpreter tool when it harms performance, adjust temperature from 0.7 to 0.0 for deterministic tasks, and replace multi-turn agentic strategies with single-call workflows when appropriate.

The results in Figure 4 reflect these differences. Mem0 achieves minimal cost reduction despite maintaining memory, as it cannot fundamentally change the underlying strategy’s computational structure. Dynamic Cheatsheet’s costs grow unboundedly as accumulated memory increases context length in every call. EGuR achieves superior accuracy while reducing costs by 2-111 $\times$ through complete strategy adaptation.

4.4 RQ3: Is comparative strategy evaluation important for continual improvements?

We ablate the exploration level k (strategies generated per problem) to assess the importance of comparative evaluation. Figure 5 shows results for Claude 3.7 Sonnet. EGuR-ZS (zero-shot, no memory updates) serves as a stateless baseline, EGuR-1 generates one strategy per problem and learns from absolute feedback, while EGuR-5 generates five strategies enabling comparative evaluation.

Three findings emerge. First, memory is essential: EGuR-1 consistently outperforms EGuR-ZS, showing that learning from experience helps even without comparison. Second, comparative evaluation provides substantial additional benefits: EGuR-5 outperforms EGuR-1 on most tasks, with large improvements on 3-SAT and object counting. Third, cost decreases with higher k because comparative evaluation helps discover efficient strategies faster.

4.5 RQ4: Does EGuR learn novel and useful strategies from experience?

We analyze strategies generated before and after training to identify learned adaptations. EGuR exhibits several consistent patterns: for CodeAct strategies, it learns to specify allowed libraries, include useful code

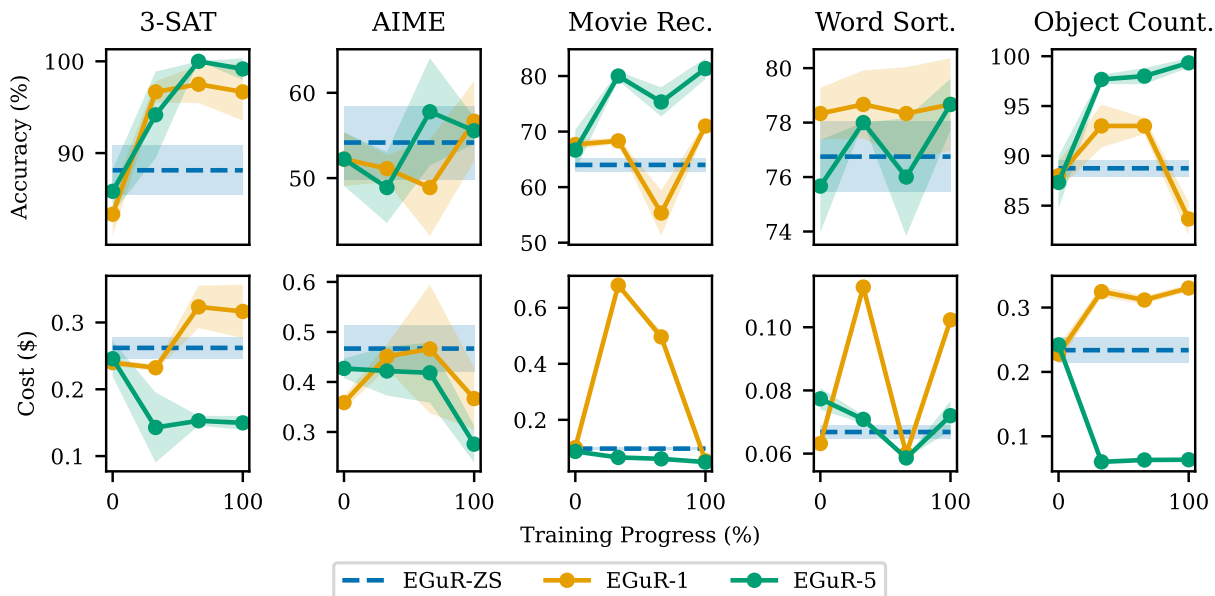


Figure 5: Ablation of exploration level in EGUR for Claude 3.7 Sonnet. Higher exploration levels (more strategies per problem) generally improve both accuracy and cost-efficiency by enabling comparative evaluation of strategy effectiveness.

snippets, and add error handling; more generally, it increases specificity when general approaches fail but simplifies when appropriate.

We also find that the code interpreter tool sometimes harms performance while increasing cost. For BBEH object counting, intuition suggests CodeAct would excel (problems involve adding many numbers), but EGUR converges on a single LLM call with detailed instructions which proves more accurate and substantially cheaper. The learned strategy includes specific guidance for parsing text, categorizing items, and handling quantity changes. The strategy is shown in Appendix D.

Similarly, for BBEH word sorting, EGUR learns to distinguish algorithmic sorting problems (which benefit from Python’s `sort`) from reasoning problems (identifying logical mistakes in explanations). For the latter, it adopts CoT with verifier-based fallback rather than CodeAct.

These findings show EGUR learns meaningful heuristics for when to use tools versus LLM-based reasoning, when to increase versus decrease computational investment, and how to tailor strategies through prompts to problem characteristics. Additional examples appear in Appendix D.

5 Related Work

Prompting strategies. Prompting methods including Chain-of-Thought (CoT) [13], Self-Consistency [11], Program-of-Thoughts (PoT) [14], and many other approaches [19–22] elicit intermediate reasoning traces but rely on fixed, non-adaptive scaffolds. More recent approaches like CodeAct [16], Self-Discover [6], and Meta-Prompting [23] dynamically compose solution solving methods per instance, but still do not leverage past experience or adapt their strategies over time. These methods are thus dynamic at inference but remain stateless across instances.

Prompt adaptation from inference-time state or training. Another line of work adapts prompts or parameters within fixed agent workflows. DSPy [2] and TextGrad [5] are systems to optimize prompts or even the inputs to fixed strategy structures using a training dataset with labels or an answer verifier, while Dynamic Cheatsheet [6] and Buffer of Thoughts [7] accumulate and retrieve artifacts during inference for instance-level adaptation. Concurrent to this work, training-free GRPO [24] was proposed as an effective

method for online prompt optimization which is similar to our GRPO inspired approach to producing effective memory updates from the Consolidator, and Agentic Context Engineering [25] improves over Dynamic Cheatsheet by providing more structure for memory updates. Systems such as DSPy are highly general and allow for many different algorithms for offline prompt optimization [26–28]. Other methods such as Per-Instance Program Synthesis (PIPS) [29], model routing [30–32], and proprietary products like GPT-5 [33] rely on trained routers to select between a fixed set of strategies on a per-instance basis. These approaches achieve stronger performance through experience but are limited in their adaptability based on only modifying the input to a fixed strategy or selecting between a fixed number of strategies.

Full strategy adaptation. Beyond prompt tuning, ADAS [8] treats agent design as a search problem over Python programs defining agents, enabling strategy-level adaptation. However, ADAS and similar approaches [34–37] operate offline and at the task level, requiring multiple episodes to converge. In contrast, our method introduces an *online adaptive meta-agent* that performs dynamic, instance-level adaptation and test-time strategy search, combining the benefits of prior work on dynamic strategies, memory-based learners, and automated design in a unified framework.

6 Conclusion

We have described a system, EGUR, which generates complete strategies—computational procedures involving LLM calls, tools, and control logic—dynamically at inference time based on accumulated experience. Specifically, a Guide component generates problem-specific strategy specifications conditioned on past experience, while a Consolidator maintains structured memory of strategy effectiveness. Unlike existing adaptive systems that use memory for textual steering of an LLM or agent, our approach has much greater flexibility, enabling adaptation of all strategy components (prompts, sampling parameters, tool configurations, and control logic). Across five benchmarks, EGUR achieves up to 14% relative accuracy improvements while reducing computational costs by up to 111x, with both metrics improving as the system gains experience. Our analysis reveals EGUR learns meaningful heuristics: when to deploy expensive agentic strategies versus lightweight workflows, when tools help versus harm, and how to tailor computational investment to problem difficulty.

Our approach has some limitations that suggest directions for future work. First, EGUR uses ground truth feedback from verifiers to learn effectively. Exploring whether this feedback can be replaced with weaker signals—such as LLM-based evaluation—is an important direction. Second, the effectiveness of the system hinges on the Guide’s zero-shot strategy generation capabilities, which may be suboptimal for unfamiliar problem types. In this case, it may be necessary to train or optimize the Guide through reinforcement learning or other training methods. Similarly, the Consolidator relies on an LLM to manage memory, which may not optimally balance memory size against information utility, and may benefit from meta-learning approaches for more effective memory management.

References

- [1] Lianmin Zheng, Liangsheng Yin, Zhiqiang Xie, Chuyue Livia Sun, Jeff Huang, Cody Hao Yu, Shiyi Cao, Christos Kozyrakis, Ion Stoica, Joseph E Gonzalez, et al. Sglang: Efficient execution of structured language model programs. *Advances in neural information processing systems*, 37:62557–62583, 2024.
- [2] Omar Khattab, Arnav Singhvi, Paridhi Maheshwari, Zhiyuan Zhang, Keshav Santhanam, Sri Vardhamanan A, Saiful Haq, Ashutosh Sharma, Thomas T. Joshi, Hanna Moazam, Heather Miller, Matei Zaharia, and Christopher Potts. DSPy: Compiling declarative language model calls into state-of-the-art pipelines. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=sY5N0zY50d>.
- [3] Yichen Huang and Lin F Yang. Gemini 2.5 pro capable of winning gold at imo 2025. *arXiv preprint arXiv:2507.15855*, 2025.

- [4] Gheorghe Comanici, Eric Bieber, Mike Schaekermann, Ice Pasupat, Noveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*, 2025.
- [5] Mert Yuksekgonul, Federico Bianchi, Joseph Boen, Sheng Liu, Zhi Huang, Carlos Guestrin, and James Zou. Textgrad: Automatic “differentiation” via text. *arXiv preprint arXiv:2406.07496*, 2024.
- [6] Mirac Suzgun, Mert Yuksekgonul, Federico Bianchi, Dan Jurafsky, and James Zou. Dynamic cheat-sheet: Test-time learning with adaptive memory. *arXiv preprint arXiv:2504.07952*, 2025.
- [7] Ling Yang, Zhaochen Yu, Tianjun Zhang, Shiyi Cao, Minkai Xu, Wentao Zhang, Joseph E Gonzalez, and Bin Cui. Buffer of thoughts: Thought-augmented reasoning with large language models. *Advances in Neural Information Processing Systems*, 37:113519–113544, 2024.
- [8] Shengran Hu, Cong Lu, and Jeff Clune. Automated design of agentic systems. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=t9U3LW7JVX>.
- [9] Prateek Chhikara, Dev Khant, Saket Aryan, Taranjeet Singh, and Deshraj Yadav. Mem0: Building production-ready ai agents with scalable long-term memory. *arXiv preprint arXiv:2504.19413*, 2025.
- [10] Ching-An Cheng, Allen Nie, and Adith Swaminathan. Trace is the next autodiff: Generative optimization with rich feedback, execution traces, and llms. *Advances in Neural Information Processing Systems*, 37:71596–71642, 2024.
- [11] Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc V Le, Ed H. Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. Self-consistency improves chain of thought reasoning in language models. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=1PL1NIMMrw>.
- [12] Erik Schluntz and Barry Zhang. Building effective ai agents. <https://www.anthropic.com/engineering/building-effective-agents>, December 2024. Accessed: 2025-09-30.
- [13] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837, 2022.
- [14] Wenhui Chen, Xueguang Ma, Xinyi Wang, and William W. Cohen. Program of thoughts prompting: Disentangling computation from reasoning for numerical reasoning tasks. *Transactions on Machine Learning Research*, 2023. ISSN 2835-8856. URL <https://openreview.net/forum?id=YfZ4ZPt8zd>.
- [15] Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, et al. Self-refine: Iterative refinement with self-feedback. *Advances in Neural Information Processing Systems*, 36:46534–46594, 2023.
- [16] Xingyao Wang, Yangyi Chen, Lifan Yuan, Yizhe Zhang, Yunzhu Li, Hao Peng, and Heng Ji. Executable code actions elicit better llm agents. In *Forty-first International Conference on Machine Learning*, 2024.
- [17] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Yang Wu, et al. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*, 2024.
- [18] Mehran Kazemi, Bahare Fatemi, Hritik Bansal, John Palowitch, Chrysovalantis Anastasiou, Saniket Vaibhav Mehta, Lalit K Jain, Virginia Aglietti, Disha Jindal, Peter Chen, et al. Big-bench extra hard. *arXiv preprint arXiv:2502.19187*, 2025.
- [19] Qing Lyu, Shreya Havaldar, Adam Stein, Li Zhang, Delip Rao, Eric Wong, Marianna Apidianaki, and Chris Callison-Burch. Faithful chain-of-thought reasoning. In *The 13th International Joint Conference on Natural Language Processing and the 3rd Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics (IJCNLP-AACL 2023)*, 2023.

- [20] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik R Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. In *The Eleventh International Conference on Learning Representations*, 2023. URL https://openreview.net/forum?id=WE_vluYUL-X.
- [21] Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc V Le, and Ed H. Chi. Least-to-most prompting enables complex reasoning in large language models. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=WZH7099tgfM>.
- [22] Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Tom Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models. *Advances in neural information processing systems*, 36:11809–11822, 2023.
- [23] Mirac Suzgun and Adam Tauman Kalai. Meta-prompting: Enhancing language models with task-agnostic scaffolding. *arXiv preprint arXiv:2401.12954*, 2024.
- [24] Yuzheng Cai, Siqi Cai, Yuchen Shi, Zihan Xu, Lichao Chen, Yulei Qin, Xiaoyu Tan, Gang Li, Zongyi Li, Haojia Lin, et al. Training-free group relative policy optimization. *arXiv preprint arXiv:2510.08191*, 2025.
- [25] Qizheng Zhang, Changran Hu, Shubhangi Upasani, Boyuan Ma, Fenglu Hong, Vamsidhar Kamanuru, Jay Rainton, Chen Wu, Mengmeng Ji, Hanchen Li, et al. Agentic context engineering: Evolving contexts for self-improving language models. *arXiv preprint arXiv:2510.04618*, 2025.
- [26] Lakshya A Agrawal, Shangyin Tan, Dilara Soylu, Noah Ziemis, Rishi Khare, Krista Opsahl-Ong, Arnav Singhvi, Herumb Shandilya, Michael J Ryan, Meng Jiang, et al. Gepa: Reflective prompt evolution can outperform reinforcement learning. *arXiv preprint arXiv:2507.19457*, 2025.
- [27] 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*, 2024. URL <https://openreview.net/forum?id=Bb4VGOWELI>.
- [28] Yongchao Zhou, Andrei Ioan Muresanu, Ziwen Han, Keiran Paster, Silviu Pitis, Harris Chan, and Jimmy Ba. Large language models are human-level prompt engineers. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=92gvk82DE->.
- [29] Adam Stein, Neelay Velingker, Mayur Naik, and Eric Wong. Once upon an input: Reasoning via per-instance program synthesis. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2025. URL <https://openreview.net/forum?id=XsQNqRdcdh>.
- [30] Lingjiao Chen, Matei Zaharia, and James Zou. FrugalGPT: How to use large language models while reducing cost and improving performance. *Transactions on Machine Learning Research*, 2024. ISSN 2835-8856. URL <https://openreview.net/forum?id=cSimKw5p6R>.
- [31] Ruochen Wang, Sohyun An, Minhao Cheng, Tianyi Zhou, Sung Ju Hwang, and Cho-Jui Hsieh. One prompt is not enough: Automated construction of a mixture-of-expert prompts. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=edHNL40DWu>.
- [32] Isaac Ong, Amjad Almahairi, Vincent Wu, Wei-Lin Chiang, Tianhao Wu, Joseph E. Gonzalez, M Waleed Kadous, and Ion Stoica. RouteLLM: Learning to route LLMs from preference data. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=8sSqNtaMr>.
- [33] OpenAI. Introducing gpt-5. <https://openai.com/index/introducing-gpt-5/>, Aug 2025.
- [34] Jiayi Zhang, Jinyu Xiang, Zhaoyang Yu, Fengwei Teng, Xiong-Hui Chen, Jiaqi Chen, Mingchen Zhuge, Xin Cheng, Sirui Hong, Jinlin Wang, Bingnan Zheng, Bang Liu, Yuyu Luo, and Chenglin Wu. AFlow: Automating agentic workflow generation. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=z5uVAKwmjf>.

Table 2: Taxonomy of related work by **S@T**=Stateful at test, **Adapted**=What is adapted in response to the state, input, or training samples, **Training-free**=If the method does not require an initial training stage

Method	S@T	Adapted	Training-free
Stateless strategies (human designed)			
CoT			✓
Self-Consistency			✓
PoT (non-iterative)			✓
CodeAct [16]			✓
Adapts prompts from state/training			
Self-Discover [6]		Prompt	✓
DSPy [2]		Prompts	
TextGrad [5]		Prompts, Inputs	
Buffer of Thoughts [7]	✓	Prompt	✓
Dynamic Cheatsheet [6]	✓	Prompt	✓
Full strategy adaptation			
FlowReasoner [38]		Strategy	
ADAS [8]		Strategy	
EGUR	✓	Strategy	✓

- [35] Mingchen Zhuge, Wenyi Wang, Louis Kirsch, Francesco Faccio, Dmitrii Khizbullin, and Jürgen Schmidhuber. GPTSwarm: Language agents as optimizable graphs. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=uTC9AFXIhg>.
- [36] Zelong Li, Shuyuan Xu, Kai Mei, Wenyue Hua, Balaji Rama, Om Raheja, Hao Wang, He Zhu, and Yongfeng Zhang. Autoflow: Automated workflow generation for large language model agents. *arXiv preprint arXiv:2407.12821*, 2024.
- [37] Guibin Zhang, Kaijie Chen, Guancheng Wan, Heng Chang, Hong Cheng, Kun Wang, Shuyue Hu, and Lei Bai. Evoflow: Evolving diverse agentic workflows on the fly. *arXiv preprint arXiv:2502.07373*, 2025.
- [38] Hongcheng Gao, Yue Liu, Yufei He, Longxu Dou, Chao Du, Zhijie Deng, Bryan Hooi, Min Lin, and Tianyu Pang. Flowreasoner: Reinforcing query-level meta-agents. *arXiv preprint arXiv:2504.15257*, 2025.

A Related Work Taxonomy

A taxonomy of the related work is shown in Table 2.

B Strategies

B.1 Syntax

The syntax of strategies is separated into values, expressions and programs.

Values

$$\begin{array}{lcl}
 v & ::= & \text{str} \mid \text{float} \mid \text{true} \mid \text{false} \mid \text{null} \\
 & & \mid [v \text{ (, } v) *] \\
 & & \mid \{\text{str} : v \text{ (, str : } v) * \}
 \end{array}$$

Values are the typical Python values including strings, floats, booleans, the None value (written null) as well as lists and dictionaries.

Expressions

$$e ::= x \mid v \mid e[e] \mid e ++ e \mid e \oplus e \mid e + e \mid e - e \mid e * e \mid e \div e \\ \mid \text{lambda } x. e \mid e(e)$$

Expressions consist of either a variable x , a value, or various operations of expressions including list indexing, list append $++$, dictionary append $\oplus$, float arithmetic $(+, -, *, \div)$, lambda functions, and function application.

Programs

$$p ::= \text{base_proc} \mid \text{return} \mid \text{pure } e \mid \text{get} \mid \text{put } e \\ \mid p; p \\ \mid p \parallel p \\ \mid \text{if } p \text{ then } p \text{ else } p \\ \mid \text{fix } p. p$$

Programs are then either a base process, a return statement (which directly returns the input), an expression lifted to a program using **pure**, a **get** which returns the current state of the program, a **put** of an expression which updates the state with the expression, or a sequential or parallel composition of programs, a conditional of programs, or a fixpoint of a program.

B.2 Semantics

We now provide a denotational semantics for the strategy language. First, a program in this language is a *stateful process* mapping between two sets:

$$\text{Proc}(A, B, \Sigma) : (A \times \Sigma) \rightarrow (B \times \Sigma)$$

where Σ is the state space of the process. Intuitively, a process can be thought of as a mapping from A to B which has access to some state Σ . To simplify how we compose processes, we use the state monad:

$$\text{State}(A) : \Sigma \rightarrow (A \times \Sigma)$$

and then rewrite the process type as $\text{Proc}(A, B, \Sigma) : A \rightarrow \text{State}(B)$ which can be thought of as a function which maps an input in A to a computation which when given a state will produce an output in B and an updated state. The choice of the state space Σ determines how processes share information when composed as well as allows for tracking information such as execution traces, cost, and strategy structure. For instance, we use $\Sigma = M \times P \times \text{Trace} \times \text{Cost}$ where M is the current conversation log of any model calls, P is the execution state for any code execution, Trace is the current execution trace, and Cost is the current execution cost of the strategy.

The program meaning function $\llbracket \cdot \rrbracket$ is now a mapping from the syntactic forms defined above to processes.

$$\llbracket \cdot \rrbracket : \text{Program} \longrightarrow \text{Proc}(A, B, \Sigma)$$

$$\begin{aligned}
\llbracket \text{base_proc} \rrbracket &= \text{base_proc}, \\
\llbracket \text{return} \rrbracket &= \lambda a. \lambda c. (a, c), \\
\llbracket \text{pure } f \rrbracket &= \lambda a. \lambda c. (\llbracket f \rrbracket(a), c), \\
\llbracket \text{get} \rrbracket &= \lambda a. \lambda c. (c, c), \\
\llbracket \text{put } e \rrbracket &= \lambda a. \lambda c. (a, \llbracket e \rrbracket), \\
\llbracket p_1; p_2 \rrbracket &= \lambda a. \lambda c. \text{let } (b, c_1) \leftarrow \llbracket p_1 \rrbracket(a)(c) \\
&\quad \text{in } \llbracket p_2 \rrbracket(b)(c_1), \\
\llbracket p_1 \parallel p_2 \rrbracket(a)(c) &= \lambda a. \lambda c. \text{let } (b_1, c_1) \leftarrow \llbracket p_1 \rrbracket(a)(c) \text{ and } (b_2, c_2) \leftarrow \llbracket p_2 \rrbracket(a)(c), \\
&\quad \text{in } ((b_1, b_2), (c_1, c_2)), \\
\llbracket \text{if } p_1 \text{ then } p_2 \text{ else } p_3 \rrbracket(s) &= \lambda a. \lambda c. \text{let } (b, c_1) \leftarrow \llbracket p_1 \rrbracket(a)(c) \\
&\quad \text{in } \begin{cases} \llbracket p_2 \rrbracket(a, c_1), & \text{if } b = \text{true} \\ \llbracket p_3 \rrbracket(a, c_1), & \text{if } b = \text{false} \end{cases} \\
\llbracket \text{fix } p_1. p_2 \rrbracket &= \text{lfp}_{\sqsubseteq} (p \mapsto \llbracket p_2 \rrbracket_{p_1 \mapsto p}).
\end{aligned}$$

B.3 Guide and Consolidator

The Guide is a strategy which takes a query as input and outputs a strategy for solving the problem, as well as has access to a state which it can use to store information between queries. Therefore, the Guide is just a higher-order strategy, or a stateful mapping from inputs to strategies and its type is the following:

$$\text{Guide}(A, B, \Sigma) : \text{Proc}(A, \text{Proc}(A, B, \sigma), \Sigma).$$

The memory updater can also be written as a strategy which takes the current experience, the current running state, and outputs the updated state:

$$\text{Consolidator}(\text{Exp}, \Sigma) : \text{Proc}(\text{Exp}, \Sigma, \Sigma).$$

B.4 Strategy Cost

We define the cost of a strategy recursively based on the costs of the base processes $P \in \mathcal{P}$ defined by $c_P(s) \in \mathbb{R}_{\geq 0}$, the cost of running P on input s .

$$\begin{aligned}
\text{Cost} : \text{Term} \times S &\longrightarrow \mathbb{R}_{\geq 0} \cup \{\infty\}, \\
\text{Cost}(t, s) &= \begin{cases} c_P(s), & \text{if } t = P \in \mathcal{P}, \\ \text{Cost}(t_1, s) + \mathbb{E}_{s' \sim \llbracket t_1 \rrbracket(s)} [\text{Cost}(t_2, s')], & \text{if } t = t_1; t_2, \\ \text{Cost}(t_1, s) + \text{Cost}(t_2, s), & \text{if } t = t_1 \parallel t_2, \\ \begin{cases} \text{Cost}(t_1, s), & b(s) = \text{true}, \\ \text{Cost}(t_2, s), & b(s) = \text{false}, \end{cases} & \text{if } t = \text{if } b \text{ then } t_1 \text{ else } t_2, \\ \text{Cost}(t[X \mapsto \text{fix } X. t], s), & \text{if } t = \text{fix } X. t \text{ (least fixed point)}. \end{cases}
\end{aligned}$$

B.5 EGUR Implemented as a Strategy

We show how EGUR is represented as a strategy in Algorithm 2 where we make sequential and parallel composition of processes explicit.

Algorithm 2 Experience-Guided Reasoner

Require: Question q , context Σ , and exploration factor k .

Ensure: Answer a and updated context Σ' .

```

1: update (fun ( $q, \Sigma$ ):  $\Sigma + \{\text{"question"} : q\}$ );           // store the question in context
2: Guide $k$ ;           // generate  $k$  strategies conditioned on the question and context
3: (recfun ExecuteStrategies:           // execute all strategies in parallel
4:   if IsEmpty
5:   then return
6:   else (GetFirst; Trace; Verify) || (PopFirst; ExecuteStrategies)
7: );
8: update (fun ( $e, \Sigma$ ):  $\Sigma + \{\text{"answer"} : e[0][0]\}$ );   // store the answer from the first strategy
9: Consolidator           // update the context from the  $k$  experiences
10: get "answer"

```

<pre> 1 CallLLM </pre>	<pre> 1 (CallLLM CallLLM 2 CallLLM CallLLM 3 CallLLM); 4 MajorityVote </pre>	<pre> 1 CallLLM; 2 (if ContainsCode 3 then ExecCode 4 else return) </pre>	<pre> 1 recfun Eval-Opt : 2 CallOptLLM; 3 (if EvalLLM 4 then return 5 else Eval-Opt) </pre>
(a) CoT	(b) Self-Consistency	(c) Code	(d) Eval-Opt

Figure 6: Implementation of CoT, Self-Consistency, Code, and Eval-Opt as strategies

B.6 Common Strategies

In this section, we describe the five strategies shown in Table 1. We also show the implementations of the strategies in our strategy language in Figure 6.

The CoT strategy is a single LLM inference call without any tools, the Code strategy is a single round of code writing and executing the code with an interpreter to produce the answer, the Eval-Opt strategy is an evaluator-optimizer loop where the optimizer produces candidate solutions and the evaluator critiques the solutions, the Self-Consistency strategy samples 10 solutions from the CoT strategy and then returns the majority answer, and the CodeAct strategy uses multiple rounds of code generation and execution to solve the problem, and is prototypical of an *agentic* workflow.

C Full results

The prequential evaluation results for all models is included in Table 3.

Full results for all methods on all three models and five holdout datasets across four evenly spaced checkpoints during training are shown in Table 4, Table 5, Table 6, Table 7, and Table 8.

D Example Generated Strategies

The following shows an example of a strategy generated with GPT-OSS-120B to solve a sarcasm task (Sarc Triples) from BBEH. We can see that this strategy includes a prompt tailored to the specific problem, majority voting, and modification of the temperature parameter to the called model.

Generated Strategy

```

1 def strategy(x):
2     """
3     Detect sarcasm in three Reddit replies.

```

Table 3: Prequential Accuracy (%) and Cost (\$) on each dataset. EGUR consistently achieves the best accuracy-cost trade-offs, with bold indicating best accuracy and underline indicating second-best. Ours- k denotes EGUR with exploration level k .

Model	Method	3-SAT		AIME		Movie Rec.		Word Sort		Object Count	
		Acc	Cost	Acc	Cost	Acc	Cost	Acc	Cost	Acc	Cost
Claude	CodeAct	0.770	<u>0.257</u>	0.611	<u>0.478</u>	0.570	<u>0.137</u>	0.770	0.067	0.800	<u>0.098</u>
	CodeAct+Mem0	0.860	<u>0.325</u>	<u>0.589</u>	0.548	0.540	0.141	0.760	0.085	<u>0.870</u>	0.174
	DynamicCheatsheet	<u>0.899</u>	76.353	0.400	8.615	0.840	9.152	0.700	8.231	0.870	24.895
	EGuR-5	0.960	0.152	0.589	0.292	<u>0.700</u>	0.067	0.800	<u>0.073</u>	0.960	0.075
Qwen3	CodeAct	0.875	<u>0.032</u>	0.622	<u>0.010</u>	<u>0.680</u>	<u>0.002</u>	0.700	<u>0.003</u>	<u>0.400</u>	<u>0.017</u>
	CodeAct+Mem0	<u>0.875</u>	0.037	<u>0.656</u>	0.011	0.750	0.003	<u>0.680</u>	0.004	0.360	0.019
	DynamicCheatsheet	0.125	0.570	0.467	0.124	0.550	0.205	0.450	0.171	0.180	0.469
	EGuR-3	0.695	0.005	0.767	0.006	0.650	0.002	0.650	0.001	0.540	0.005
GPT	CodeAct	0.655	0.236	<u>0.778</u>	0.073	0.550	0.091	0.780	0.087	<u>0.690</u>	0.047
	CodeAct+Mem0	<u>0.640</u>	<u>0.219</u>	0.800	<u>0.078</u>	<u>0.530</u>	0.061	<u>0.780</u>	<u>0.129</u>	0.790	<u>0.059</u>
	EGuR-3	0.650	0.100	0.756	0.095	0.350	<u>0.064</u>	0.580	0.091	0.620	0.090

Table 4: Accuracy (%) across training progress for different methods on 3-SAT and AIME datasets. Values shown as mean $\pm$ standard deviation across seeds. Claude uses EGuR-5, GPT and Qwen use EGuR-3. **Bold** indicates best performance within each model, underlined indicates second best within each model.

Model	Method	3-SAT				AIME			
		0%	33%	66%	100%	0%	33%	66%	100%
Claude-3.5-Sonnet	CodeAct	79.4 $\pm$ 4.1	79.4 $\pm$ 4.1	79.4 $\pm$ 4.1	79.4 $\pm$ 4.1	48.3 $\pm$ 5.0	48.3 $\pm$ 5.0	48.3 $\pm$ 5.0	48.3 $\pm$ 5.0
	EGuR-ZS	88.1$\pm$2.7	88.1$\pm$2.7	88.1$\pm$2.7	88.1$\pm$2.7	54.2$\pm$4.3	54.2$\pm$4.3	54.2$\pm$4.3	54.2$\pm$4.3
	Mem0	74.2 $\pm$ 2.4	79.2 $\pm$ 6.6	78.3 $\pm$ 4.2	86.7 $\pm$ 5.1	53.3 $\pm$ 0.0	58.9$\pm$6.8	55.6 $\pm$ 1.6	46.7 $\pm$ 8.2
	DC	64.1 $\pm$ 12.1	72.3 $\pm$ 8.8	73.3 $\pm$ 2.7	78.8 $\pm$ 8.0	41.1 $\pm$ 3.1	38.9 $\pm$ 4.2	<u>28.9$\pm$4.2</u>	34.4 $\pm$ 4.2
	EGuR-5	<u>85.8$\pm$2.4</u>	94.2$\pm$4.7	100.0$\pm$0.0	99.2$\pm$1.2	52.2 $\pm$ 3.1	48.9 $\pm$ 4.2	57.8$\pm$6.3	55.6$\pm$1.6
GPT-OSS-120B	CodeAct	88.8 $\pm$ 5.7	88.8$\pm$5.7	88.8 $\pm$ 5.7	88.8$\pm$5.7	79.6$\pm$6.8	79.6$\pm$6.8	79.6$\pm$6.8	79.6$\pm$6.8
	EGuR-ZS	83.1 $\pm$ 2.7	83.1 $\pm$ 2.7	83.1 $\pm$ 2.7	83.1 $\pm$ 2.7	65.8 $\pm$ 3.6	65.8 $\pm$ 3.6	65.8 $\pm$ 3.6	65.8 $\pm$ 3.6
	Mem0	91.7$\pm$1.2	86.7 $\pm$ 5.9	91.7$\pm$2.4	<u>84.2$\pm$2.4</u>	<u>76.6$\pm$3.4</u>	60.0 $\pm$ 6.7	69.1 $\pm$ 9.1	<u>79.4$\pm$16.1</u>
	DC	10.8 $\pm$ 6.6	6.7 $\pm$ 2.4	19.2 $\pm$ 4.2	<u>13.3$\pm$4.2</u>	<u>45.6$\pm$4.2</u>	44.4 $\pm$ 5.7	53.3 $\pm$ 7.2	<u>44.4$\pm$4.2</u>
	EGuR-3	85.8 $\pm$ 2.4	80.8 $\pm$ 9.6	35.0 $\pm$ 6.1	69.2 $\pm$ 12.3	65.6 $\pm$ 8.7	<u>75.6$\pm$4.2</u>	<u>73.3$\pm$7.2</u>	76.7 $\pm$ 4.7
Qwen3-Next-80B	CodeAct	63.8 $\pm$ 7.4	63.8 $\pm$ 7.4	63.8 $\pm$ 7.4	63.8 $\pm$ 7.4	<u>74.2$\pm$4.9</u>	74.2 $\pm$ 4.9	74.2 $\pm$ 4.9	74.2 $\pm$ 4.9
	EGuR-ZS	80.0$\pm$4.7	<u>80.0$\pm$4.7</u>	80.0$\pm$4.7	80.0$\pm$4.7	78.8$\pm$7.5	78.8$\pm$7.5	78.8$\pm$7.5	78.8$\pm$7.5
	Mem0	75.0 $\pm$ 0.0	<u>52.5$\pm$0.0</u>	57.5 $\pm$ 0.0	<u>67.5$\pm$0.0</u>	66.7 $\pm$ 0.0	76.7 $\pm$ 0.0	80.0$\pm$0.0	76.7 $\pm$ 0.0
	EGuR-3	72.5 $\pm$ 0.0	90.0$\pm$0.0	<u>75.0$\pm$0.0</u>	45.0 $\pm$ 0.0	56.7 $\pm$ 0.0	72.4 $\pm$ 0.0	60.0 $\pm$ 0.0	55.2 $\pm$ 0.0

```

4  Repeatedly query the model, extract the list after "FINAL ANSWER:",
5  and return the majority-voted label string.
6  """
7  import re
8  from collections import Counter
9
10 # Prompt that enforces step-by-step reasoning and a strict final answer format
11 base_prompt = (
12     "You are given three (post, reply) pairs. For each reply decide whether it is "
13     "sarcastic (1) or not sarcastic (0). Think carefully about each pair, then output "
14     "the three labels in order as a comma-separated list. Write the list after a line "
15     "that starts exactly with 'FINAL ANSWER:'. \n"
16     f"Problem: {x}"
17 )
18
19 collected_answers = []
20
21 # Number of independent attempts (adjustable; 5 balances cost and reliability)
22 for _ in range(5):
23     # First generation
24     out = M(base_prompt, temperature=0.5, max_tokens=2000)[0]
25
26     # Try to extract the answer
27     ans = extract_final_answer(out)
28     if not ans:
29         # Ask the model to repeat with the required format

```

Table 5: Accuracy (%) across training progress for different methods on Movie Recommendation, Word Sorting, and Object Counting datasets. Values shown as mean $\pm$ standard deviation across seeds. Claude uses EGuR-5, GPT and Qwen use EGuR-3. **Bold** indicates best performance within each model, underlined indicates second best within each model.

Model	Method	Movie Rec.				Word Sort.				Object Count.			
		0%	33%	66%	100%	0%	33%	66%	100%	0%	33%	66%	100%
Claude-3.5-Sonnet	CodeAct	57.2 $\pm$ 1.1	57.2 $\pm$ 1.1	57.2 $\pm$ 1.1	57.2 $\pm$ 1.1	75.8 $\pm$ 0.4	75.8 $\pm$ 0.4	75.8 $\pm$ 0.4	75.8 $\pm$ 0.4	85.2 $\pm$ 4.8	85.2 $\pm$ 4.8	85.2 $\pm$ 4.8	85.2 $\pm$ 4.8
	EGuR-ZS	<u>64.0$\pm$1.2</u>	64.0 $\pm$ 1.2	64.0 $\pm$ 1.2	64.0 $\pm$ 1.2	76.8$\pm$1.3	<u>76.8$\pm$1.3</u>	76.8$\pm$1.3	<u>76.8$\pm$1.3</u>	88.8$\pm$0.8	<u>88.8$\pm$0.8</u>	<u>88.8$\pm$0.8</u>	88.8 $\pm$ 0.8
	Mem0	53.7 $\pm$ 4.6	56.0 $\pm$ 1.4	55.7 $\pm$ 6.6	51.0 $\pm$ 2.9	<u>76.7$\pm$0.9</u>	74.7 $\pm$ 1.7	73.7 $\pm$ 2.5	74.3 $\pm$ 1.2	86.3 $\pm$ 0.9	86.3 $\pm$ 2.5	85.7 $\pm$ 3.3	85.3 $\pm$ 1.7
	DC	59.7 $\pm$ 3.1	80.6$\pm$2.0	79.0$\pm$2.5	<u>77.7$\pm$2.1</u>	64.8 $\pm$ 3.4	68.8 $\pm$ 2.5	71.5 $\pm$ 0.6	68.3 $\pm$ 0.9	70.7 $\pm$ 3.3	85.3 $\pm$ 1.5	88.7 $\pm$ 1.7	<u>90.1$\pm$2.3</u>
	EGuR-5	66.7$\pm$3.7	<u>80.0$\pm$0.8</u>	<u>75.3$\pm$2.6</u>	81.3$\pm$1.9	75.7 $\pm$ 1.7	78.0$\pm$0.0	<u>76.0$\pm$2.2</u>	78.7$\pm$0.9	<u>87.3$\pm$2.6</u>	97.7$\pm$0.5	98.0$\pm$0.8	99.3$\pm$0.5
GPT-OSS-120B	CodeAct	72.8 $\pm$ 3.3	<u>72.8$\pm$3.3</u>	<u>72.8$\pm$3.3</u>	72.8 $\pm$ 3.3	66.0$\pm$1.7	66.0$\pm$1.7	66.0$\pm$1.7	66.0$\pm$1.7	39.5 $\pm$ 7.1	39.5 $\pm$ 7.1	39.5 $\pm$ 7.1	39.5$\pm$7.1
	EGuR-ZS	75.0$\pm$3.2	75.0$\pm$3.2	75.0$\pm$3.2	75.0$\pm$3.2	<u>64.8$\pm$4.0</u>	<u>64.8$\pm$4.0</u>	<u>64.8$\pm$4.0</u>	<u>64.8$\pm$4.0</u>	29.5 $\pm$ 1.8	29.5 $\pm$ 1.8	29.5 $\pm$ 1.8	29.5 $\pm$ 1.8
	Mem0	73.3 $\pm$ 0.9	72.0 $\pm$ 2.9	72.0 $\pm$ 1.6	71.0 $\pm$ 2.4	64.0 $\pm$ 3.6	63.3 $\pm$ 2.4	63.7 $\pm$ 2.1	63.7 $\pm$ 1.7	35.3 $\pm$ 2.5	<u>41.0$\pm$5.4</u>	<u>41.0$\pm$5.4</u>	38.0 $\pm$ 2.2
	DC	<u>73.7$\pm$1.9</u>	57.0 $\pm$ 1.4	52.3 $\pm$ 0.5	59.7 $\pm$ 0.5	47.3 $\pm$ 0.5	37.7 $\pm$ 4.5	45.0 $\pm$ 2.4	39.3 $\pm$ 2.5	13.0 $\pm$ 4.5	15.7 $\pm$ 1.7	21.7 $\pm$ 2.1	8.7 $\pm$ 5.2
	EGuR-3	62.3 $\pm$ 3.4	60.3 $\pm$ 1.7	62.3 $\pm$ 3.1	<u>74.7$\pm$2.5</u>	63.3 $\pm$ 2.6	63.0 $\pm$ 2.8	<u>65.0$\pm$1.6</u>	64.0 $\pm$ 3.6	<u>37.3$\pm$2.6</u>	48.3$\pm$4.8	57.3$\pm$1.2	<u>39.0$\pm$4.5</u>
Qwen3-Next-80B	CodeAct	<u>46.5$\pm$1.5</u>	46.5$\pm$1.5	<u>46.5$\pm$1.5</u>	46.5$\pm$1.5	<u>75.5$\pm$2.1</u>	75.5$\pm$2.1	<u>75.5$\pm$2.1</u>	<u>75.5$\pm$2.1</u>	75.5$\pm$3.2	75.5$\pm$3.2	75.5$\pm$3.2	<u>75.5$\pm$3.2</u>
	EGuR-ZS	39.2 $\pm$ 1.7	39.2 $\pm$ 1.7	39.2 $\pm$ 1.7	39.2 $\pm$ 1.7	64.4 $\pm$ 2.3	64.4 $\pm$ 2.3	64.4 $\pm$ 2.3	64.4 $\pm$ 2.3	<u>74.9$\pm$4.5</u>	<u>74.9$\pm$4.5</u>	<u>74.9$\pm$4.5</u>	74.9 $\pm$ 4.5
	Mem0	47.0$\pm$0.0	<u>45.0$\pm$0.0</u>	47.0$\pm$0.0	<u>43.0$\pm$0.0</u>	77.0$\pm$0.0	75.0$\pm$0.0	77.0$\pm$0.0	78.0$\pm$0.0	74.0 $\pm$ 0.0	63.0 $\pm$ 0.0	70.0 $\pm$ 0.0	77.0 $\pm$ 0.0
	EGuR-3	41.0 $\pm$ 0.0	32.0 $\pm$ 0.0	28.0 $\pm$ 0.0	40.0 $\pm$ 0.0	65.0 $\pm$ 0.0	43.0 $\pm$ 0.0	68.0 $\pm$ 0.0	64.0 $\pm$ 0.0	60.0 $\pm$ 0.0	44.0 $\pm$ 0.0	38.0 $\pm$ 0.0	37.0 $\pm$ 0.0

Table 6: Cost (\$) across training progress for different methods on 3-SAT and AIME datasets. Values shown as mean $\pm$ standard deviation across seeds. Claude uses EGuR-5, GPT and Qwen use EGuR-3. **Bold** indicates lowest cost within each model, underlined indicates second lowest within each model.

Model	Method	3-SAT				AIME			
		0%	33%	66%	100%	0%	33%	66%	100%
Claude-3.5-Sonnet	CodeAct	0.398 $\pm$ 0.049	0.398 $\pm$ 0.049	0.398 $\pm$ 0.049	0.398 $\pm$ 0.049	0.814 $\pm$ 0.112	0.814 $\pm$ 0.112	0.814 $\pm$ 0.112	0.814 $\pm$ 0.112
	EGuR-ZS	<u>0.262$\pm$0.016</u>	<u>0.262$\pm$0.016</u>	<u>0.262$\pm$0.016</u>	<u>0.262$\pm$0.016</u>	<u>0.467$\pm$0.046</u>	<u>0.467$\pm$0.046</u>	<u>0.467$\pm$0.046</u>	<u>0.467$\pm$0.046</u>
	Mem0	0.386 $\pm$ 0.029	0.371 $\pm$ 0.040	0.349 $\pm$ 0.059	0.396 $\pm$ 0.011	0.880 $\pm$ 0.075	0.718 $\pm$ 0.107	0.837 $\pm$ 0.091	0.733 $\pm$ 0.124
	DC	8.969 $\pm$ 1.020	10.056 $\pm$ 1.431	10.572 $\pm$ 1.095	9.951 $\pm$ 1.944	1.133 $\pm$ 0.149	1.371 $\pm$ 0.395	2.146 $\pm$ 0.109	2.259 $\pm$ 0.280
	EGuR-5	0.246$\pm$0.026	0.143$\pm$0.052	0.153$\pm$0.007	0.150$\pm$0.010	0.427$\pm$0.020	0.422$\pm$0.049	0.418$\pm$0.060	0.276$\pm$0.037
GPT-OSS-120B	CodeAct	0.036 $\pm$ 0.007	0.036 $\pm$ 0.007	0.036 $\pm$ 0.007	0.036 $\pm$ 0.007	<u>0.004$\pm$0.000</u>	0.004$\pm$0.000	<u>0.004$\pm$0.000</u>	<u>0.004$\pm$0.000</u>
	EGuR-ZS	<u>0.005$\pm$0.001</u>	<u>0.005$\pm$0.001</u>	<u>0.005$\pm$0.001</u>	<u>0.005$\pm$0.001</u>	<u>0.004$\pm$0.001</u>	<u>0.004$\pm$0.001</u>	<u>0.004$\pm$0.001</u>	<u>0.004$\pm$0.001</u>
	Mem0	0.027 $\pm$ 0.002	0.031 $\pm$ 0.005	0.038 $\pm$ 0.009	0.035 $\pm$ 0.004	0.006 $\pm$ 0.002	0.010 $\pm$ 0.001	0.007 $\pm$ 0.001	0.011 $\pm$ 0.003
	DC	0.019 $\pm$ 0.002	0.014 $\pm$ 0.001	0.017 $\pm$ 0.001	0.021 $\pm$ 0.002	0.018 $\pm$ 0.000	0.017 $\pm$ 0.001	0.028 $\pm$ 0.001	0.032 $\pm$ 0.001
	EGuR-3	0.004$\pm$0.000	0.003$\pm$0.000	0.004$\pm$0.001	0.004$\pm$0.001	0.003$\pm$0.000	0.006 $\pm$ 0.002	0.004$\pm$0.001	0.003$\pm$0.001
Qwen3-Next-80B	CodeAct	0.251 $\pm$ 0.024	0.251 $\pm$ 0.024	0.251 $\pm$ 0.024	0.251 $\pm$ 0.024	<u>0.104$\pm$0.015</u>	<u>0.104$\pm$0.015</u>	0.104 $\pm$ 0.015	0.104 $\pm$ 0.015
	EGuR-ZS	<u>0.137$\pm$0.013</u>	<u>0.137$\pm$0.013</u>	<u>0.137$\pm$0.013</u>	<u>0.137$\pm$0.013</u>	0.120 $\pm$ 0.023	0.120 $\pm$ 0.023	0.120 $\pm$ 0.023	0.120 $\pm$ 0.023
	Mem0	0.212 $\pm$ 0.000	0.208 $\pm$ 0.000	0.223 $\pm$ 0.000	0.228 $\pm$ 0.000	0.105 $\pm$ 0.000	0.115 $\pm$ 0.000	<u>0.080$\pm$0.000</u>	<u>0.089$\pm$0.000</u>
	EGuR-3	0.089$\pm$0.000	0.097$\pm$0.000	0.135$\pm$0.000	0.096$\pm$0.000	0.057$\pm$0.000	0.038$\pm$0.000	0.041$\pm$0.000	0.054$\pm$0.000

```

30 clarification = (
31     "Please respond again, making sure to include a line that starts exactly with "
32     "'FINAL ANSWER:' followed by the three labels separated by commas."
33 )
34 out = M(clarification + "\nPrevious output:\n" + out,
35         temperature=0, max_tokens=2000)[0]
36 ans = extract_final_answer(out)
37
38 if ans:
39     # Normalise whitespace and store
40     collected_answers.append(ans.strip())
41
42 if not collected_answers:
43     return "No answer found."
44
45 # Majority vote over the collected answer strings
46 most_common = Counter(collected_answers).most_common(1)[0][0]
47 return most_common

```

For Claude 3.7 Sonnet, we show generated strategies before and after training on the BBEH Object Counting task. The first strategy below, from before training answers the particular problem it was generated for incorrectly and cost \$0.166 while the following strategy answers the same problem correctly and only costs \$0.062.

Table 7: Cost (\$) across training progress for different methods on Movie Recommendation and Word Sorting datasets. Values shown as mean $\pm$ standard deviation across seeds. Claude uses EGuR-5, GPT and Qwen use EGuR-3. **Bold** indicates lowest cost within each model, underlined indicates second lowest within each model.

Model	Method	Movie Rec.				Word Sort.			
		0%	33%	66%	100%	0%	33%	66%	100%
Claude-3.5-Sonnet	CodeAct	0.142 $\pm$ 0.004	0.142 $\pm$ 0.004	0.142 $\pm$ 0.004	0.142 $\pm$ 0.004	0.060$\pm$0.007	0.060$\pm$0.007	<u>0.060$\pm$0.007</u>	0.060$\pm$0.007
	EGuR-ZS	<u>0.097$\pm$0.005</u>	<u>0.097$\pm$0.005</u>	<u>0.097$\pm$0.005</u>	<u>0.097$\pm$0.005</u>	<u>0.067$\pm$0.002</u>	<u>0.067$\pm$0.002</u>	0.067 $\pm$ 0.002	<u>0.067$\pm$0.002</u>
	Mem0	0.161 $\pm$ 0.004	0.156 $\pm$ 0.005	0.166 $\pm$ 0.007	0.166 $\pm$ 0.005	0.082 $\pm$ 0.005	0.084 $\pm$ 0.004	0.085 $\pm$ 0.006	0.078 $\pm$ 0.001
	DC	0.826 $\pm$ 0.011	2.525 $\pm$ 0.047	2.710 $\pm$ 0.034	2.885 $\pm$ 0.014	1.214 $\pm$ 0.047	2.822 $\pm$ 0.133	5.426 $\pm$ 0.066	4.322 $\pm$ 0.068
	EGuR-5	0.087$\pm$0.001	0.066$\pm$0.002	0.061$\pm$0.001	0.049$\pm$0.001	0.077 $\pm$ 0.003	0.071 $\pm$ 0.001	0.059$\pm$0.000	0.072 $\pm$ 0.005
GPT-OSS-120B	CodeAct	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	<u>0.003$\pm$0.000</u>	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000
	EGuR-ZS	<u>0.002$\pm$0.000</u>	<u>0.002$\pm$0.000</u>	<u>0.002$\pm$0.000</u>	0.002$\pm$0.000	0.001$\pm$0.000	0.001$\pm$0.000	0.001$\pm$0.000	<u>0.001$\pm$0.000</u>
	Mem0	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	0.003 $\pm$ 0.000	0.004 $\pm$ 0.000	0.004 $\pm$ 0.000	0.004 $\pm$ 0.000	0.004 $\pm$ 0.000
	DC	0.038 $\pm$ 0.001	0.093 $\pm$ 0.001	0.099 $\pm$ 0.000	0.069 $\pm$ 0.001	0.038 $\pm$ 0.001	0.059 $\pm$ 0.001	0.082 $\pm$ 0.002	0.052 $\pm$ 0.001
	EGuR-3	0.001$\pm$0.000	0.001$\pm$0.000	0.001$\pm$0.000	0.005 $\pm$ 0.000	<u>0.001$\pm$0.000</u>	<u>0.002$\pm$0.000</u>	<u>0.001$\pm$0.000</u>	0.001$\pm$0.000
Qwen3-Next-80B	CodeAct	0.096 $\pm$ 0.003	0.096 $\pm$ 0.003	0.096 $\pm$ 0.003	0.096 $\pm$ 0.003	0.083 $\pm$ 0.003	0.083 $\pm$ 0.003	0.083 $\pm$ 0.003	0.083 $\pm$ 0.003
	EGuR-ZS	0.073 $\pm$ 0.005	0.073 $\pm$ 0.005	0.073 $\pm$ 0.005	0.073 $\pm$ 0.005	<u>0.078$\pm$0.009</u>	<u>0.078$\pm$0.009</u>	<u>0.078$\pm$0.009</u>	0.078$\pm$0.009
	Mem0	<u>0.071$\pm$0.000</u>	<u>0.066$\pm$0.000</u>	<u>0.067$\pm$0.000</u>	<u>0.063$\pm$0.000</u>	0.100 $\pm$ 0.000	0.118 $\pm$ 0.000	0.102 $\pm$ 0.000	0.110 $\pm$ 0.000
	EGuR-3	0.047$\pm$0.000	0.049$\pm$0.000	0.043$\pm$0.000	0.056$\pm$0.000	0.068$\pm$0.000	0.052$\pm$0.000	0.053$\pm$0.000	<u>0.083$\pm$0.000</u>

Table 8: Cost (\$) across training progress for different methods on Object Counting dataset. Values shown as mean $\pm$ standard deviation across seeds. Claude uses EGuR-5, GPT and Qwen use EGuR-3. **Bold** indicates lowest cost within each model, underlined indicates second lowest within each model.

Model	Method	Object Count.			
		0%	33%	66%	100%
Claude-3.5-Sonnet	CodeAct	0.095$\pm$0.005	0.095 $\pm$ 0.005	0.095 $\pm$ 0.005	0.095 $\pm$ 0.005
	EGuR-ZS	0.234 $\pm$ 0.020	0.234 $\pm$ 0.020	0.234 $\pm$ 0.020	0.234 $\pm$ 0.020
	Mem0	0.125 $\pm$ 0.001	0.127 $\pm$ 0.005	0.125 $\pm$ 0.005	0.128 $\pm$ 0.009
	DC	1.539 $\pm$ 0.046	8.155 $\pm$ 0.295	8.803 $\pm$ 0.931	7.156 $\pm$ 0.397
	EGuR-5	0.242 $\pm$ 0.003	0.060$\pm$0.001	0.063$\pm$0.001	0.064$\pm$0.001
GPT-OSS-120B	CodeAct	0.020 $\pm$ 0.001	0.020 $\pm$ 0.001	0.020 $\pm$ 0.001	0.020 $\pm$ 0.001
	EGuR-ZS	<u>0.006$\pm$0.001</u>	<u>0.006$\pm$0.001</u>	0.006$\pm$0.001	<u>0.006$\pm$0.001</u>
	Mem0	0.020 $\pm$ 0.003	0.022 $\pm$ 0.001	0.021 $\pm$ 0.003	0.019 $\pm$ 0.000
	DC	0.200 $\pm$ 0.009	0.250 $\pm$ 0.019	0.194 $\pm$ 0.010	0.128 $\pm$ 0.048
	EGuR-3	0.005$\pm$0.000	0.004$\pm$0.000	<u>0.006$\pm$0.000</u>	0.006$\pm$0.001
Qwen3-Next-80B	CodeAct	<u>0.049$\pm$0.001</u>	<u>0.049$\pm$0.001</u>	<u>0.049$\pm$0.001</u>	<u>0.049$\pm$0.001</u>
	EGuR-ZS	0.086 $\pm$ 0.010	0.086 $\pm$ 0.010	0.086 $\pm$ 0.010	0.086 $\pm$ 0.010
	Mem0	0.055 $\pm$ 0.000	0.057 $\pm$ 0.000	0.064 $\pm$ 0.000	0.065 $\pm$ 0.000
	EGuR-3	0.049$\pm$0.000	0.032$\pm$0.000	0.029$\pm$0.000	0.025$\pm$0.000

BBEH Object Counting strategy from Claude 3.7 Sonnet before training

```

1  async def strategy(x):
2      SYSTEM_PROMPT = """You are an expert at solving counting and categorization
3  problems with Python code.
4
5  Task: Solve the question by iteratively writing and executing Python code.
6
7  Rules:
8  - Reply with code inside a single ```python ...``` block. The code will be executed
9  for you; anything you print will be returned to you next round.
10 - Read the execution result, and then decide on how to proceed next. You must iterate
11 until you have solved the problem.
12 - When ready, return the solution exactly as: FINAL ANSWER: <answer>
13 (You may print this from code or write it directly outside the code block.)
14
15 Constraints:
16 - Code must be self-contained, use only the Python standard library, and avoid
17 network/file I/O and pip installs.
18
19 Approach:
20 1. Parse the text to identify all items the narrator ("I") has
21 2. Categorize each item as either a musical instrument or an animal/insect
22 3. Count the total number of each category
23 4. Calculate the absolute difference between the two totals"""
24
25  PARSE_ERROR_PROMPT = "Parsing error: I couldn't find a ```python ... ``` code \
26  block or a line starting with 'FINAL ANSWER:' in the model output. \

```

```

27 Please either:\n1) Provide valid Python within ```python ... ``` OR\n2) \
28 Reply directly with a line exactly: FINAL ANSWER: <answer> with no trailing text."
29
30 max_rounds = 20
31 messages = [{"role": "system", "content": SYSTEM_PROMPT}]
32
33 messages.append({"role": "user", "content": f"Problem: {x}"})
34 out = await model_call(messages, temperature=0.7, thinking=True, max_tokens=20000)
35 messages.append({"role": "assistant", "content": out})
36 execution_state = {}
37 for i in range(max_rounds):
38     # First check for any output code to execute
39     blocks = extract_python_code(out)
40     if blocks:
41         block = "\n\n".join(blocks)
42         exec_res, execution_state = python_exec(block, execution_state)
43
44         # Check for final answer in execution output
45         m = extract_final_answer(exec_res)
46         if m:
47             return m.strip()
48
49         # No final answer, so take another turn of model output
50         messages.append({
51             "role": "user",
52             "content": f"The output of executing the generated Python code \
53                 is:\n{exec_res}"})
54         out = await model_call(
55             messages, temperature=0.7, thinking=True, max_tokens=20000)
56         messages.append({"role": "assistant", "content": out})
57         continue
58
59         # If no code, then check for a final answer
60         m = extract_final_answer(out)
61         if m:
62             return m.strip()
63         messages.append({"role": "user", "content": PARSE_ERROR_PROMPT})
64         out = await model_call(messages, temperature=0.7, thinking=True, max_tokens=20000)
65         messages.append({"role": "assistant", "content": out})
66
67 messages.append({
68     "role": "user",
69     "content": "Now output the final answer on a single line as \
70         'FINAL ANSWER: answer' with no trailing text."})
71 out = await model_call(messages, temperature=0.7, thinking=True, max_tokens=20000)
72 m = extract_final_answer(out)
73 if m:
74     return m.strip()
75 else:
76     raise Exception("No answer found.")

```

BBEH Object Counting strategy from Claude 3.7 Sonnet after training

```

1 async def strategy(x):
2     PROMPT = """Solve the following problem step by step:
3
4     1. First, carefully read through the text and identify all items that "I"
5       (the narrator) have. Ignore items that belong to other people
6       (grandmother, friend, sister, brother, father, mother, aunt, grandfather,
7       cousin, uncle).
8     2. For each item that belongs to "I", categorize it as either:
9       - An animal/insect (e.g., ducks, rabbits, walrus, hens, moose, wasps, bees,
10        chameleons, newts, vultures, otter, tigers, shrimps, elks, baboons, yaks,
11        boa snakes, anteaters, magpies, fish, narwhals, chimpanzees, antelopes)
12       - A musical instrument (e.g., harps, panpipes, balafons, marimbas,
13        electric guitars, recorders, clavichord, electric pianos, tanpuras, ouds,

```

```

14     bass guitars)
15 - Other items (cars like nissan pathfinder, kia niro, honda accord,
16   nissan frintier, toyota highlander, hyundai elantra, mazda2, buick enclave,
17   hyundai santa fe) which we will ignore for this problem
18 3. Create a table with columns: Item, Category (Animal/Insect, Musical Instrument,
19   Other), Quantity.
20 4. For items with stories about changing quantities, ALWAYS use the FINAL quantity
21   mentioned. For example:
22   - "initially I had X, then I lost Y, but then I got Z, but then I lost W,
23     so at the end I am left with V" → Count as V
24   - "initially I had X, then I also got Y making it a total of Z, then I got W
25     making it a total of A, but then I lost B of them so eventually I am left
26     with C" → Count as C
27 5. Count the total number of animals/insects that "I" have.
28 6. Count the total number of musical instruments that "I" have.
29 7. Calculate the absolute difference between these two totals.
30
31 Be methodical and organize your work clearly. After completing your analysis,
32 write the final answer on a single line as 'FINAL ANSWER: <answer>' with no trailing
33 text.
34
35 Problem: {x}"""
36
37 MISSING_ANSWER_FEEDBACK = "Your last response did not include a line exactly of \
38   the form 'FINAL ANSWER: <answer>'. Please try again to output a final answer,\
39   responding with exactly one line 'FINAL ANSWER: <answer>'."
40
41 messages = [{"role": "user", "content": PROMPT.format(x=x)}]
42 for _ in range(5):
43     output = await model_call(
44         messages, temperature=0.7, thinking=True, max_tokens=20000)
45     # Find FINAL ANSWER
46     ans = extract_final_answer(output)
47     if ans:
48         return ans
49     messages.append({"role": "assistant", "content": output})
50     messages.append({"role": "user", "content": MISSING_ANSWER_FEEDBACK})
51
52 # One last attempt with a more direct prompt
53 messages.append({
54     "role": "user",
55     "content": "Based on your analysis, what is the absolute difference between \
56     the total number of animals/insects and musical instruments that the narrator \
57     has? Provide only the numerical answer in the format 'FINAL ANSWER: <number>'."})
58 output = await model_call(
59     messages, temperature=0.7, thinking=True, max_tokens=20000)
60 ans = extract_final_answer(output)
61 if ans:
62     return ans
63 else:
64     raise Exception("No answer found after multiple attempts")

```

E Example Learned Context

The final memory resulting from EGUR-5 with Claude 3.7 Sonnet on BBEH Movie Recommendations is shown below. The memory consists of several memory entries and a single entry containing the (task, best strategy) pair which is used for solving future problems.

Learned Context

```

1 <memory_entry-1>
2 When analyzing similarity between items for group appeal (movies, books, etc.), consider these key
  ↳ factors:
3 1. Era/time period consistency (items from the same era tend to have more similar appeal)

```

```

4 2. Language/accessibility (foreign language items create significant outliers)
5 3. Target demographics (mainstream vs. niche appeal)
6 4. Critical reception and popular appeal
7 5. Production values and stylistic approaches
8
9 For movie similarity specifically, the most reliable indicators of similar group appeal are:
10 - Same language and era
11 - Similar accessibility factors
12 - Comparable mainstream vs. niche positioning
13 - Consistent tone and pacing
14 </memory_entry-1>
15
16 <memory_entry-2>
17 For similarity analysis problems, the Chain of Thought (CoT) approach with deep thinking tends to be more
18   ↳ cost-effective than more complex strategies like CodeAct or Evaluator-Optimizer when the task involves
19   ↳ subjective judgment rather than algorithmic calculation.
20
21 The CoT approach with a well-structured prompt that guides the analysis through specific criteria (genres,
22   ↳ audience, reception, etc.) and asks for explicit rating/ranking of options provides sufficient
23   ↳ structure for the model to make accurate comparisons without requiring the overhead of code execution
24   ↳ or multiple model calls.
25 </memory_entry-2>
26
27 [...]
28
29 <memory_entry-5>
30 Task: Analyzing which group of movies has the most similar audience appeal
31 Best Strategy:
32 <strategy>
33 async def strategy(x):
34     PROMPT = f"""You are an expert film critic tasked with analyzing which group of movies would have the
35       ↳ most similar appeal to the same audience. Think step by step and provide your final answer on a
36       ↳ single line as 'FINAL ANSWER: <option letter>' with no trailing text.
37
38 When analyzing similarity between movies for group appeal, consider these key factors:
39 1. Era/time period consistency (movies from the same era tend to have more similar appeal)
40 2. Language/accessibility (foreign language films create significant outliers)
41 3. Target demographics (mainstream vs. niche appeal)
42 4. Critical reception and popular appeal
43 5. Production values and stylistic approaches
44 6. Psychological appeal patterns (puzzle-solving vs. emotional engagement vs. escapism)
45 7. Viewing experience requirements (active vs. passive engagement)
46 8. Narrative complexity and accessibility
47 9. Emotional responses elicited (fear, humor, intellectual stimulation, etc.)
48
49 For movie similarity specifically, the most reliable indicators of similar group appeal are:
50 - Same language and era
51 - Similar accessibility factors
52 - Comparable mainstream vs. niche positioning
53 - Consistent tone and pacing
54
55 Significant outliers that reduce group cohesion include:
56 - Animated shorts among feature films
57 - Family/children's content mixed with adult-oriented films
58 - Critically panned films grouped with acclaimed ones
59 - Comedic/light entertainment mixed with serious dramas
60 - Foreign language films mixed with English-language films
61 - Adult/erotic content mixed with general audience films
62
63 For each option:
64 1. Analyze each movie in the group individually
65 2. Compare each movie against every other movie in the group
66 3. Identify specific outliers that would reduce group cohesion
67 4. Rate the overall cohesion on a 1-10 scale
68 5. Explain your reasoning thoroughly
69
70 Problem: {x}

```

```

65 Think through each option carefully and systematically.
66 """
67     MISSING_ANSWER_FEEDBACK = "Your last response did not include a line exactly of the form 'FINAL
    ↳ ANSWER: <option letter>'. Please try again to output a final answer, responding with exactly one
    ↳ line 'FINAL ANSWER: <option letter>'."
68
69     messages = [{"role": "user", "content": PROMPT}]
70     for _ in range(5):
71         output = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000)
72         # Find FINAL ANSWER
73         ans = extract_final_answer(output)
74         if ans:
75             return ans
76         messages.append({"role": "assistant", "content": output})
77         messages.append({"role": "user", "content": MISSING_ANSWER_FEEDBACK})
78     raise Exception("No answer found")
79 </strategy>
80 </memory_entry-5>

```

F Prompts

We include the complete prompts for the Guide and Consolidator in the following sections. These prompts define the two strategies defining EGUR.

F.1 Guide

The prompt for the Guide is included below.

Guide Prompt

```

1 You are designing a reasoning strategy, not solving the target problem.
2
3 A reasoning strategy is a Python function that describes the design of an LLM workflow/agent to solve
    ↳ a particular problem.
4 It should not directly contain or compute the final answer to the target problem. Instead, it should
    ↳ provide the workflow/environment for an LLM system to arrive at the correct answer.
5 The strategy doesn't need to describe the explicit steps to take to solve the problem and can instead just
    ↳ set up an agent which can figure out how to solve the problem itself (see the CodeAct example below).
6 Once you output a runnable strategy, you will be given the reasoning trace from the strategy execution and
    ↳ you can optionally update any saved memory.
7
8 Definition:
9 A reasoning strategy is an async function strategy(x) where:
10 - x (string) is the problem statement.
11 - The function can contain calls to the following functions:
12   - model_call: (list[dict], temperature, thinking, max_tokens) -> Awaitable[str]: an async call to a
    ↳ model which takes conversation history (list of dict) and some keyword generation parameters as
    ↳ input and outputs the response. You may use asyncio.gather to perform concurrent calls to the
    ↳ model.
13   - extract_python_code: str -> Optional[List[str]]: a parsing function to extract all code from
    ↳ Python markdown blocks using the regex r````(?:python|py)\s*(.*?)\s*```.
14   - extract_final_answer: str -> Optional[str]: a parsing function to extract the answer for text in
    ↳ the format of "FINAL ANSWER: answer".
15   - python_exec: (str, dict) -> (str, dict): a function which takes the code string and state dict as
    ↳ input and returns the output of executing the code as well as the updated state dict. This can
    ↳ only execute python code using standard libraries, so it does not support any imports which must
    ↳ be installed with pip.
16 - The function can use global variables MEMORY_ENTRY_K which are set to the string containing the entire
    ↳ body of <memory_entry-k></memory_entry-k>.
17 - The strategy function must return the final answer as a string.
18
19 Important:

```

```

20 - **Do not** attempt to answer the problem yourself or write code without LLM calls which solves the
    ↪ problem.
21 - **Do not** implement the `python_exec` function or the `extract_final_answer` function as it will be
    ↪ provided if you need it.
22 - Your output must define the `strategy` Python function in a <strategy></strategy> block**.
23 - The function must be executable without changes and cannot import any non-standard libraries.
24 - If including code in a prompt, be careful about including quotes within a string. Use different types of
    ↪ quotes for the outer and inner quotes or explicitly escape all the inner quotes.
25 - The strategy should return the final answer only with no other text.
26 - You should reference any available memory below to help design the best strategy.
27
28 **Strategy motifs**
29 Some commonly used elements which can be composed together to create the reasoning strategy:
30 - chain of thought (by asking to think step by step)
31 - code interpreter (using python_exec)
32 - iterative feedback to improve answers or correct any possible runtime errors (feeding the output from
    ↪ the model back to itself in case of error)
33 - majority voting (sample multiple times with temperature and take the most common answer)
34 - deep thinking mode (enabled with the thinking=True flag to `model_call` which automatically sets
    ↪ temperature to 1.0. Temperature MUST BE greater than 0 (preferably 0.6-1.0) when thinking=True.)
35 - problem decomposition (first split to subproblems, then solve each subproblem either with a separate LLM
    ↪ call or all at once)
36 - evaluator and optimizer (keep trying to solve until the proposed answer is deemed correct by a
    ↪ verification method which can be implemented as a tool call or LLM call)
37
38 **Example implementations of some strategies:**
39 <example>
40 CoT with deep thinking (a single prompt to an LLM to solve the problem with retries if answer is missing)
41 <strategy>
42 ```python
43 async def strategy(x):
44     PROMPT = "Solve the following problem step by step and write the final answer on a single line as
    ↪     'FINAL ANSWER: <answer>' with no trailing text:\nProblem: {x}"
45     MISSING_ANSWER_FEEDBACK = "Your last response did not include a line exactly of the form 'FINAL
    ↪     ANSWER: <answer>'. Please try again to output a final answer, responding with exactly one line
    ↪     'FINAL ANSWER: <answer>'."
46
47     messages = [{"role": "user", "content": PROMPT.format(x=x)}]
48     for _ in range(10):
49         output = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000) # Note that
    ↪         you can set thinking=False to turn off deep thinking for simpler problems
50         # Find FINAL ANSWER
51         ans = extract_final_answer(output)
52         if ans:
53             return ans
54         messages.append({"role": "assistant", "content": output})
55         messages.append({"role": "user", "content": MISSING_ANSWER_FEEDBACK})
56     raise Exception("No answer found")
57 ```
58 </strategy>
59 </example>
60
61 <example>
62 CodeAct (loop of writing code and executing the code until answer is reached)
63 <strategy>
64 ```python
65 async def strategy(x):
66     SYSTEM_PROMPT = """You are an expert at solving problems with Python code.
67
68     Task: Solve the question by iteratively writing and executing Python code.
69
70     Rules
71     - Reply with code inside a single ```python ...``` block. The code will be executed for you; anything you
    ↪     print will be returned to you next round.
72     - Read the execution result, and then decide on how to procede next. You must iterate until you have
    ↪     solved the problem.
73     - When ready, return the solution exactly as: FINAL ANSWER: <answer>

```

```

74 (You may print this from code or write it directly outside the code block.)
75
76 Constraints
77 - Code must be self-contained, use only the Python standard library, and avoid network/file I/O and pip
  ↳ installs."""
78 PARSE_ERROR_PROMPT = "Parsing error: I couldn't find a ```python ... ``` code block or a line starting
  ↳ with 'FINAL ANSWER:' in the model output. Please either:\n1) Provide valid Python within ```python
  ↳ ... ``` OR\n2) Reply directly with a line exactly: FINAL ANSWER: <answer> with no trailing text."
79
80 max_rounds = 20
81 messages = [{"role": "system", "content": SYSTEM_PROMPT}]
82
83 messages.append({"role": "user", "content": f"Problem: {x}"})
84 out = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000)
85 messages.append({"role": "assistant", "content": out})
86 execution_state = {}
87 for i in range(max_rounds):
88     # First check for any output code to execute
89     # We must always first execute any code or else we may parse out a "FINAL ANSWER" print statement
  ↳ before executing it
90     blocks = extract_python_code(out)
91     if blocks:
92         block = "\n\n".join(blocks)
93         exec_res, execution_state = python_exec(block, execution_state)
94
95         # Check for final answer in execution output
96         m = extract_final_answer(exec_res)
97         if m:
98             return m.strip()
99
100         # No final answer, so take another turn of model output
101         messages.append({"role": "user", "content": f"The output of executing the generated Python
  ↳ code is:\n{exec_res}"})
102         out = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000)
103         messages.append({"role": "assistant", "content": out})
104         continue
105
106         # If no code, then check for a final answer
107         m = extract_final_answer(out)
108         if m:
109             return m.strip()
110         messages.append({"role": "user", "content": PARSE_ERROR_PROMPT})
111         out = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000)
112         messages.append({"role": "assistant", "content": out})
113
114 messages.append({"role": "user", "content": "Now output the final answer on a single line as 'FINAL
  ↳ ANSWER: answer' with no trailing text."})
115 out = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000)
116 m = extract_final_answer(out)
117 if m:
118     return m.strip()
119 else:
120     raise Exception("No answer found.")
121 ...
122 </strategy>
123 </example>
124
125 <example>
126 Parallelization (sample multiple times with temperature and then take the most common response. For very
  ↳ ambiguous or challenging problems, sampling 20+ times can be helpful):
127 <strategy>
128 async def strategy(x):
129     PROMPT = "Solve the problem and output the final answer on a single line as 'FINAL ANSWER: <answer>'
  ↳ with no trailing text.\nProblem: {x}"
130
131     outputs = []
132     messages = [{"role": "user", "content": PROMPT.format(x=x)}]
133     async def get_response():

```

```

134     # use higher temperature for more variation
135     out = await model_call(messages, temperature=1.0, thinking=True, max_tokens=20000)
136
137     # Find FINAL ANSWER
138     ans = extract_final_answer(out)
139     if ans:
140         return ans
141     else:
142         return "No answer found"
143
144     # multiple independent calls can run concurrently with asyncio.gather
145     outputs = await asyncio.gather(*[get_response() for _ in range(10)])
146
147     # return the majority of the answers
148     return max(set(outputs), key = outputs.count)
149 </strategy>
150 </example>
151
152 <example>
153 Evaluator-Optimizer (one LLM call generates a solution and another provides evaluation and feedback for
    ↳ the generator to improve its solution)
154 <strategy>
155 ```python
156 async def strategy(x):
157     GENERATOR_PROMPT = "You are an expert problem solver. Solve the given problem step by step and provide
    ↳ your final answer on a single line as 'FINAL ANSWER: <answer>' with no trailing text."
158     EVALUATOR_PROMPT = "Evaluate whether the following answer to a problem is correct.\nProblem:
    ↳ {problem}\nProposed Answer: {solution}\nIf the answer is correct, respond with exactly 'FINAL
    ↳ ANSWER: CORRECT'. If the answer is incorrect, respond with 'FINAL ANSWER: <explanation of what is
    ↳ wrong and guidance for improvement>'."
159     MISSING_ANSWER_PROMPT = "The last output is missing a line of the form 'FINAL ANSWER: ...' which needs
    ↳ to be parsed out. Please try again to output a line following this format."
160
161     max_iterations = 10
162
163     # Generate initial solution
164     generator_messages = [
165         {"role": "system", "content": GENERATOR_PROMPT},
166         {"role": "user", "content": f"Problem: {x}"}
167     ]
168
169     for i in range(max_iterations):
170         # get answer from generator
171         for _ in range(5):
172             solution = await model_call(generator_messages, temperature=1.0, thinking=True,
    ↳ max_tokens=20000)
173             generator_messages.append({"role": "assistant", "content": solution})
174
175             # Check if we have a final answer
176             ans = extract_final_answer(solution)
177             if ans:
178                 break
179             else:
180                 generator_messages.append({"role": "user", "content": MISSING_ANSWER_PROMPT})
181
182         # get evaluation of the current answer
183         eval_messages = [{"role": "user", "content": EVALUATOR_PROMPT.format(problem=x, solution=ans)}]
184         for _ in range(5):
185             evaluation = await model_call(eval_messages, temperature=1.0, thinking=True, max_tokens=20000)
186
187             eval_ans = extract_final_answer(evaluation)
188             if eval_ans:
189                 break
190             else:
191                 eval_messages = [{"role": "assistant", "content": evaluation}]
192                 eval_messages.append({"role": "user", "content": MISSING_ANSWER_PROMPT})
193
194         # only return the answer if the evaluator says it is correct

```

```

195     if "CORRECT" in eval_ans:
196         return ans
197
198     # otherwise, make the generator improve the answer based on the evaluator output
199     generator_messages.append({"role": "user", "content": f"An evaluator gave the following feedback
    ↳ on your answer: {evaluation}. Please revise your answer and output the new answer on a line
    ↳ as 'FINAL ANSWER: answer'."})
200
201     # Final attempt to extract answer
202     solution = await model_call(generator_messages, temperature=1.0, thinking=True, max_tokens=20000)
203     ans = extract_final_answer(solution)
204     if ans:
205         return ans.strip()
206     else:
207         raise Exception("No answer found.")
208     ...
209 </strategy>
210 </example>
211
212 If parsing anything from the output of the model, be sure to exactly specify the format the model is
    ↳ supposed to respond in and regenerate using any failed parsing feedback if there are parsing errors.
213 It's best to keep the max_tokens around 20000 for all calls to `model_call` to avoid issues from output
    ↳ truncation, unless the memory shows that using lower max_tokens is safe.
214 Your central goal is to write a strategy that will correctly solve the problem, but if there are two
    ↳ possible strategies which are equally likely to be correct, favor the less costly strategy.
215 A strategy should never give up! The strategy must solve the provided problem even if it takes a long
    ↳ time, so the strategy should take whatever steps necessary to solve the problem correctly and be sure
    ↳ it does not result in giving up and outputting an incorrect answer. If everything goes wrong, then
    ↳ throw an exception.
216 Strategies are written at the meta-level (see how the above examples are not problem specific), and you
    ↳ should only tailor prompts to the problem if you are sure it will help.
217 Note that for many non-algorithmic tasks, a simple CoT approach (with in-context examples if available) is
    ↳ one of the best strategies provided the model is capable enough, and CodeAct is one of the best
    ↳ strategies for algorithmic tasks.
218
219 **Memory**
220 [[PAST]]
221 Use the above memory to help design the strategy. It is your job to incorporate any useful information
    ↳ from the memory into the strategy. This can be incorporated into the strategy at two levels:
222 1. The high level structure (the memory provides the best strategies for related tasks as well as other
    ↳ notes about useful strategies). Base the strategy structure on the memory as much as possible.
223 2. The prompts (e.g. include all important advice, code blocks to use, pitfalls to avoid, few-shot
    ↳ examples, etc. which the model should have in its context). To include a memory entry into a strategy
    ↳ (such as to include code into a prompt), use the global variable `MEMORY_ENTRY_K` where K is replaced
    ↳ with the id of the entry or copy the content of the memory verbatim. Include as much information from
    ↳ the notes as possible into prompts to help the models avoid pitfalls and streamline their reasoning.
224 Please mention if you are using any of the above memories when designing the strategy. Before designing a
    ↳ new strategy, check if a (task, best strategy) pair exists in the memory for the current task and use
    ↳ the best strategy if so.
225 Ignore irrelevant memory entries.
226 Overall, you want to think about how to use more computation to avoid incorrect answers when they occur
    ↳ (more verification/CodeAct iterations/voting/thinking when the problem is very hard or past problems
    ↳ were answered incorrectly) and when to use less computation (simpler strategies such as CoT when the
    ↳ problem/subproblem is simple or past problems were solved easily) to reduce cost.
227 You want to produce the strategy most likely to get a correct answer.
228
229 Target problem:
230 [[PROBLEM]]
231
232 Your Task:
233 [[TASK_PROMPT]]

```

F2 Consolidator

The prompt for the Consolidator is included below.

Consolidator Prompt

```
1 The outcomes from executing the generated strategies are shown below.
2 <experience>
3 <strategy>
4   [[PAST_STRATEGY]]
5 </strategy>
6 <trace>
7   [[STRATEGY_EXECUTION_TRACE]]
8 </trace>
9 <final_answer>[FINAL_ANSWER]</final_answer>
10 <verifier_output>[[Answer is correct! or Answer is incorrect!]]</verifier_output>
11 <cost>[[COST]]</cost>
12 </experience>
13 ...
14
15 Your task is now to update a memory to help you solve problems better in the future by generating better
16   ↳ strategies.
17
18 The memory consists of (task, best strategy) pairs where there is exactly one entry per task seen as well
19   ↳ as entries about general useful findings to improve strategies in the future.
20
21 **Workflow**
22 1. For each experience, examine the verifier output and find the strategies that resulted in the correct
23   ↳ answer. Of these correct answer strategies, find the one which has greatest reliability or lowest
24   ↳ cost. If all strategies resulted in the wrong answer, then continue to step 2 and DO NOT SAVE THE
25   ↳ STRATEGY to memory.
26 1.1. Check the current memory for existing (task, best strategy) pairs where the task matches the
27   ↳ current problem.
28 1.2. If there is an existing (task, best strategy) entry in the memory for this task, update the
29   ↳ memory entry with the strategy found in step 1 ONLY IF the strategy used above is better (in
30   ↳ reliability or lower cost) than the one already in memory.
31 2. From the reasoning traces of the strategies extract any useful information (such as code snippets and
32   ↳ useful examples) which could help in the future. DO NOT SAVE any strategy which resulted in the
33   ↳ incorrect answer.
34
35 **Extract**
36 - The best strategies that resulted in correct answers, helpful functions, heuristics, examples. Feel free
37   ↳ to save code.
38 - Unexpected behaviors of models or tools in the reasoning trace which should be remembered to help design
39   ↳ better strategies and avoid unnecessary rework or continued failures in the future
40 - When certain strategy motifs are useful or harmful
41 - Useful and working code snippets, tool configurations which can be provided in the future rather than
42   ↳ rediscovered
43 - Anything that would save significant reasoning time/cost (such as having to reimplement a code solution
44   ↳ or reducing the number of reasoning steps)
45 - Failure modes and their early detection/mitigation
46 - Effective prompts or verification methods
47 - Cost-saving shortcuts that maintain accuracy
48 - Failures in tool-use which could be fixed in the strategy
49
50 **Avoid**
51 - Problem-specific literals (numbers, names, dataset IDs) or any information which would not apply to
52   ↳ other problems
53 - Obvious insights, redundant info, incomplete code
54 - Anything that requires rework, for instance if something was derived in the reasoning trace (like a code
55   ↳ solution) then don't say "solve with code" and instead save the actual code.
56 - Inventing or coming up with anything new, for instance if the reasoning trace shows a failed solution,
57   ↳ do not fix the solution now since your new solution may also be wrong. Instead, it is better to
58   ↳ examine the reasoning trace to determine why it was wrong or what was an unsuccessful strategy to
59   ↳ avoid the mistake in the future.
60 - Saving strategies that resulted in the wrong answer
61
62 **Output**
63 To add or update the (task, best strategy) pair, output blocks of the form:
64 <memory_entry>
65 Task: [short description of the task to allow you to know which problems to use this strategy for in the
66   ↳ future. Can be at any granularity that is helpful.]
67 Best Strategy:
```

```

47 <strategy>
48 [Code for the best strategy discovered for the task]
49 </strategy>
50 </memory_entry>
51 For saving other useful information, but not full strategies, output blocks of the form:
52 <memory_entry>
53 [Actionable memory entry to improve future behavior on similar problems.]
54 </memory_entry>
55
56 To delete memory entries, output blocks of the form:
57 <delete_entry>
58 [mem-id]
59 </delete_entry>
60 For instance:
61 <delete_entry>
62 1
63 </delete_entry>
64 will delete the memory entry block of the form <memory_entry-1></memory_entry-1>.
65 When you output a new memory, DO NOT add the ID to the memory_entry block as the ID will be automatically
   ↳ assigned for you.
66 To edit a memory, delete the old memory and then output the updated memory.
67 Never delete a (task, best strategy) pair from memory.
68 Be careful about deleting entries because other memory entries may refer to the deleted memory entry or
   ↳ even use the variable `MEMORY_ENTRY_K` within code.
69
70 **Existing Memory**
71 [[past_memories]]
72 Memory length: [[memory_len]] tokens
73
74 Output nothing if you do not want to change anything.
75 Try not to remove any memories which could be useful in the future.
76 Also, ensure the memory stays shorter than 10000 tokens or else it will distract more than it helps. If
   ↳ the memory is above 10k tokens, try to only edit the memory rather than make additions.
77 Before outputting the memory edits, write a short justification for your changes.

```