

SELFGOAL: Your Language Agents Already Know How to Achieve High-level Goals

Anonymous ACL submission

Abstract

Language agents powered by large language models (LLMs) are increasingly valuable as decision-making tools in domains such as gaming and programming. However, these agents often face challenges in achieving high-level goals without detailed instructions and in adapting to environments where feedback is delayed. In this paper, we present SELFGOAL, a novel automatic approach designed to enhance agents’ capabilities to achieve high-level goals with limited human prior and environmental feedback. The core concept of SELFGOAL involves adaptively breaking down a high-level goal into a tree structure of more practical subgoals during the interaction with environments while identifying the most useful subgoals and progressively updating this structure. Experimental results demonstrate that SELFGOAL significantly enhances the performance of language agents across various tasks, including competitive, cooperative, and deferred feedback environments.

1 Introduction

The advancement of large language models (LLMs) (Brown et al., 2020; OpenAI, 2022, 2024) has enabled the construction of autonomous *language agents* (or LLM-based agents) to solve complex tasks in dynamic environments without task-specific training. In reality, these autonomous agents are often tasked with very broad, high-level goals, such as “winning the most money” or “succeeding in a competition”, whose ambiguous nature and delayed reward raise great challenges for autonomous task-solving. More importantly, it is not practical to frequently train these models to adapt to new goals and tasks (Zheng et al., 2023; Khot et al., 2023; Prasad et al., 2024). Therefore, a critical question arises: *How can we enable autonomous language agents to consistently achieve high-level goals without training?*

Previous works focus on creating two types of auxiliary guidance in the instructions for language agents to achieve high-level goals in tasks: prior task decomposition and post-hoc experience summarization. The former involves decomposing the task before acting, utilizing prior knowledge from LLMs to break down high-level goals into more tangible subgoals related to specific actions at hand (Yuan et al., 2023; Zheng et al., 2023; Singh et al., 2024; Liu et al., 2024). However, this line of work does not ground these subgoals into the environment during interaction, resulting in the loss of empirical guidance. In contrast, the latter allows agents to interact directly with environments and summarize valuable experiences from history (Madaan et al., 2023; Majumder et al., 2023; Zhao et al., 2024; Paul et al., 2024), *e.g.*, “X contributes to Y”. However, the difficulty of inducing rules from experience causes the guidance to be simple and unstructured, making it difficult to prioritize or adjust strategies effectively.

A natural solution to combine the best of both worlds is to dynamically decompose the task and its high-level goal during interaction with the environment. This approach requires an agent to build and use guidelines that vary in detail and aspect. A tree structure is ideal for this requirement, as it allows hierarchical organization, providing both broad overviews and detailed guidance as needed. However, this approach presents two major challenges: 1) Not all nodes are relevant to the current context during task execution, which requires selecting the most suited nodes to guide current actions. For example, “watch for bargains” is a more prudent choice than “bid on the most expensive item” when budget is tight; 2) The granularity of guidance provided by nodes increases with tree depth, yet the appropriate detail level varies across scenarios, making a fixed tree depth not general. For example, a generic guideline like “earn more money” is not useful in

auctions.

To tackle these challenges, we propose SELF-GOAL, a self-adaptive framework for a language agent to utilize both prior knowledge and environmental feedback to achieve high-level goals. The main idea is to build a tree of textual subgoals, where agents choose appropriate ones as the guidelines to the prompt based on the situation. Specifically, as shown in Figure 1, SELF-GOAL is featured with three main modules to operate a GOALTREE, which is constructed, updated, and utilized during task execution: 1) **Search Module** is prompted to select the top-K most suited nodes of goals based on the provided current state and existing nodes in GOALTREE, which utilizes the prior knowledge of LLMs; 2) **Decomposition Module** breaks down a goal node into a list of more concrete subgoals as subsequent leaves, ensuring an adaptive self-growth of GOALTREE. Note that we filter out the redundant nodes during decomposition based on the textual similarity between new ones and the existing nodes of goals; 3) **Act Module** takes as input the selected subgoals as guidelines, and prompts LLMs for actions for the current state. Extensive experiments in various competition and collaboration scenarios show that SELF-GOAL provides precise guidance for high-level goals and adapts to diverse environments, significantly improving language agent performance.

In summary, our contributions in this paper are as follows:

- We target the challenge of enabling autonomous language agents to consistently achieve high-level goals without the need for frequent retraining.
- We introduce SELF-GOAL, a self-adaptive framework that constructs, updates, and utilizes a GOALTREE to dynamically decompose a task’s high-level goals into subgoals during interaction with the environment.
- We conduct extensive experiments in both collaborative and competitive scenarios where agents tend to deviate from their goals. The results demonstrate that SELF-GOAL significantly enhances the capability of language agents to adhere to high-level goals consistently.

2 Related Work

Learning from Feedback Recently, LLMs have become a promising tool for building goal-directed

language agents (Huang et al., 2022a). With textual input that includes the world state, task, and interaction history, language agents are to decide the next action to achieve a goal (Lin et al., 2023; Yao et al., 2023). Several studies have explored enhancing the reasoning and planning abilities of language agents through feedback from environments. For example, Reflexion (Shinn et al., 2023) enables an agent to reflect on its failures and devise a new plan that accounts for previous mistakes. Similarly, Voyager (Wang et al., 2023a) operates in Minecraft, developing a code-based skill library from detailed feedback on its failures. Recent works (Majumder et al., 2023; Nottingham et al., 2024) analyze both failures and successes attempts, summarizing a memory of causal abstractions. However, learnings directly from feedback are often too general and not systematic, making it difficult to prioritize strategies effectively.

LLMs for Decision Making LLMs are increasingly used as policy models for decision-making in interactive environments such as robotics (Ahn et al., 2022; Huang et al., 2022b; Liu et al., 2023), textual games (Wang et al., 2023b; Zhang et al., 2024; Xie et al., 2024; Ma et al., 2024), and social tasks (Zhou et al., 2024). However, the goals in these environments, like “find a fruit” in ScienceWorld (Wang et al., 2022), are often simple and specific. For long-term, high-level goals, LLMs struggle to perform effectively (Hoang et al., 2021; Huang et al., 2019), and additional modules are needed for support (Zheng et al., 2023). In our work, we use a method that does not require updating LLM parameters, enabling language agents to consistently pursue high-level goals during interactions with environments.

Decomposition and Modularity Decomposing complex decision-making tasks into sub-tasks is a traditional method that enhances LLM task-solving capabilities (Barto and Mahadevan, 2003; Pellier et al., 2023). Approaches like Hierarchical Task Networks leverage domain knowledge, including a hand-specified library of plans, to simplify complex problems (Erol et al., 1994). Recently, some studies have assigned LLMs the role of decomposing goals. For example, Decomposed Prompting (Khot et al., 2022) uses a few-shot prompting approach to tackle multi-step reasoning tasks by breaking them into a shared library of prompts. OKR-Agent (Zheng et al., 2023) utilizes self-collaboration and self-correction mechanisms, supported by hierarchi-

cal agents, to manage task complexities. ADAPT (Prasad et al., 2024) enables LLMs to recursively re-decompose goals based on feedback in decision-making tasks. However, these approaches often decompose tasks before interaction with the environments, resulting in a lack of grounded, dynamic adjustment. To address this, we aim to combine modular goal decomposition with learning from environmental feedback.

3 Methodology

Algorithm 1: Workflow of SELFGOAL

Data: Main Goal: g_0 , Threshold: ξ , Stopping criterion

- 1 Initialize Environment state s_0 , Actor M_a , p_0 , and policy $\pi_\theta(a_i|s_i)$, $\theta = \{p_0\}$
- 2 Generate initial action-state pair $\{a_0, s_0\}$ using π_θ
- 3 Generate initial GOALTREE:
 $\mathbb{T} = g_0 \cup \text{DECOMPOSE}(g_0, \{a_0, s_0\})$
- 4 Set $t = 0$
- 5 **while** *Stopping criterion not met* **do**
- 6 $g_{i,j} = \text{SEARCH}(\mathbb{T}, s_t)$
- 7 $p_{t+1} \leftarrow \{p_t, g_{i,j}\}$
- 8 $a_{t+1} = \text{ACT}(\pi_\theta, p_{t+1}, s_t)$
- 9 $G \leftarrow \text{DECOMPOSE}(g_{i,j}, \{a_{t+1}, s_{t+1}\})$
 // Update $\mathbb{T}$
- 10 **foreach** $g \in G$ **do**
- 11 **if** $\text{cosine}(g, \mathbb{T}) < \xi$ **then**
- 12 $\mathbb{T} \leftarrow \mathbb{T} \cup g$
- 13 Increment t
- 14 **return**

When executing complex tasks with high-level goals (e.g., “forecast future stock prices”), humans usually decompose it into specific detailed subgoals (e.g., “gather historical price data and adjust predictions based on recent market events”) for effective execution (Goffaux et al., 2011). Inspired from this idea, we propose SELFGOAL in this paper, which is a non-parametric learning approach for language agents to exploit and achieve high-level goals. SELFGOAL conducts a top-down hierarchical decomposition of the high-level goal, with a tree of nodes representing useful guidance for decision-making.

In this section, we first provide an overview of how SELFGOAL works in §3.1. Next, we explain the details of three key modules (Search, Decompose and Act) in SELFGOAL that help maintain a tree of subgoals (GOALTREE) in §3.2 and guide task execution.

3.1 Overview of SELFGOAL

Problem Formulation: Tasks with High-level Goals First, we formulate the features of our studied tasks, requiring an agent to interact with a dynamic environment and evaluated based on the achievement of the high-level goal. We focus on the scenarios where an actor model M_a aims to achieve a high-level goal g_0 in an environment E through interaction. The policy employed by M_a is denoted as π_θ . At each timestep t , π_θ generates an action a_t , and the environment E returns a state s_t . This action-state pair $\{a_t, s_t\}$ is then utilized to update π_θ . Note that SELFGOAL also supports accomplishing long-horizon tasks that do not always have immediate rewards. In this case, only by completing the task M_a will be evaluated with a score according to the achievement of the goal g_0 .

Workflow of SELFGOAL SELFGOAL is a non-parametric learning algorithm for language agents, i.e., without parameter update. The workflow of SELFGOAL is shown at Algorithm 1, which models $\pi_\theta = p$ by setting p as the instruction prompt provided to M_a (an LLM), i.e., $a_t = \text{LLM}(p_t, s_{t-1})$. The policy π_θ is updated through the modifications to p , which is the modification of subgoal instructions $g_{i,j}$ (i -th layer, j -th leaf node) that best suit the current situation. Concretely, SELFGOAL is featured with two key modules, **Decomposition** and **Search** which construct and utilize a subgoal tree $\mathbb{T}$ respectively, namely GOALTREE, to interact with the environment. Setting the high-level goal of the task as the root node in GOALTREE, **Search Module** finds the nodes that are helpful for the status quo, and **Decomposition Module** decomposes the node into subgoals as leaf nodes if it is not clear enough.

3.2 Details in SELFGOAL

Search: Identifying Useful Subgoals for Current Situation In the **Search** module of SELFGOAL, we ask the backbone LLM of the agent to identify the most appropriate subgoal for the current situation, e.g., “Select K most useful sub-goals that will help you reach your main goal in the current situation ...” (see Appendix A.6 for the complete prompt). We represent the the current state s_t for timestep t as a description of the dialogue history of the interaction with the environment. We also find the leaf nodes of each branch in GOALTREE as the candidate subgoal list for LLMs to decide which ones that are useful. The LLM then selects

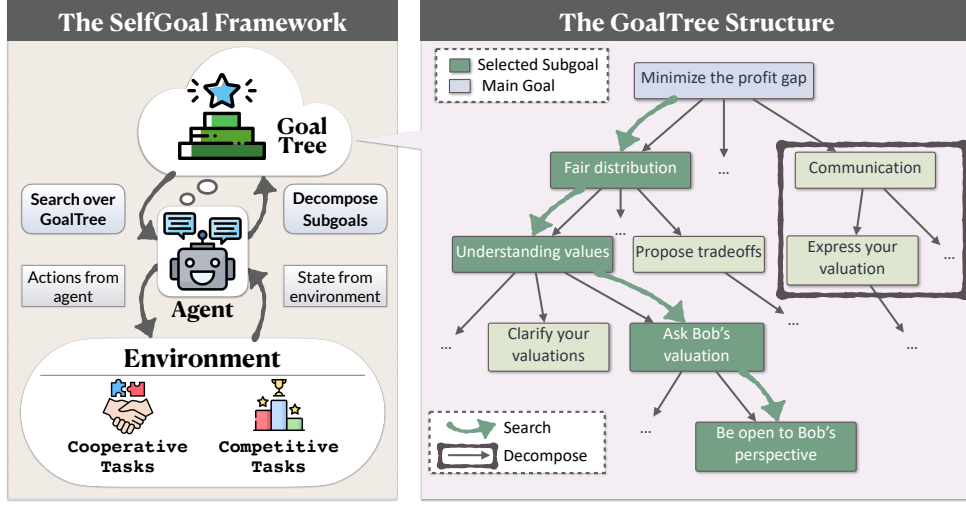


Figure 1: An overview of SELFGOAL, illustrated with a bargaining example. The agent interacts with environments, and make actions based on environmental feedback and the GOALTREE dynamically constructs, utilizes and updates with Search and Decompose Modules.

K most suitable subgoals, followed by decomposition and the update of the instruction prompt p_t at this step.

Decompose: Refine GOALTREE to Adapt to the Environment Based on the current action-state pair $\{a_t, s_t\}$, GOALTREE is updated through decomposition if it is not specific enough for useful guidance to the agent. We use the backbone LLM to break down the selected subgoal $g_{i,j}$ in the **Search Module** (initially set to g_0). We prompt the LLM with the instruction such as “What subgoals can you derive from $\{g_{i,j}\}$, based on $\{a_t, s_t\}$ ”, which generates a new set of subgoals G (see also Appendix A.6). To control the granularity of these subgoals, we apply a filtering mechanism that if the cosine similarity (Rahutomo et al., 2012) between a new subgoal and existing subgoals exceeds ξ , the current node will not be updated. Otherwise, we add the new subgoals under the current node, thus expanding the GOALTREE. Moreover, a *stopping mechanism* is designed that if no new nodes are added to the GOALTREE for N consecutive rounds, the update is stopped.

Act: Utilizing Subgoals to Take Actions After getting the subgoals from GOALTREE that are found by SELFGOAL as useful, the agent updates the instruction prompt p_t for the LLM and takes action a_t to interact with the environment. The prompt of this step can also be found in Appendix A.6.

4 Experimental Setup

4.1 Tasks and Environments

Table 1: The categorization of studied tasks.

Task	Rounds	Task Type
Public Goods Game	Single	Competitive
Guess 2/3 of the Average	Single	Cooperative
First-price Auction	Multiple	Competitive
Bargaining	Multiple	Cooperative

We evaluated SELFGOAL in four dynamic tasks with high-level goals, including **Public Goods Game**, **Guess 2/3 of the Average**, **First-price Auction**, and **Bargaining**, which are implemented by existing work (Huang et al., 2024; Chen et al., 2023; Lewis et al., 2017). As seen in Table 1, they are either single-round or multi-round games, requiring the collaboration or competition of multiple agents. Note that agents in multi-round games will only receive delayed rewards at the end of the game. In our experiments, we repeat single-round games for $T = 20$ times and multi-round games for $T = 10$ times for stable results.

Public Goods Game: GAMA-Bench We use **GAMA-Bench** (Huang et al., 2024) as the implemented environment for this game. Specifically, each of $N = 5$ players privately decides the number of tokens contributed to a public pot. The tokens in the pot are multiplied by a factor R ($1 \leq R \leq N$), and the created “public good” is distributed evenly among all players. Players keep any tokens they do not contribute. A simple calculation reveals that for each token a player contributes, their net gain is $\frac{R}{N} - 1$ (i.e., income-contribution). Since this value is

negative, it suggests that the most rational strategy for each player is to contribute no tokens. This strategy results in a Nash equilibrium (Daskalakis et al., 2009) in the game. N agents using the same backbone model and equipped with the same method (e.g., CLIN or SELFGOAL) play games with each other to observe group behavior. Following (Huang et al., 2024), we set $R = 2$.

Guess 2/3 of the Average: GAMA-Bench Using the implementation of **GAMA-Bench** (Huang et al., 2024), N players independently choose a number between 0 and 100 (Ledoux, 1981), and whoever has the number closest to two-thirds of the group’s average wins the game. This setup effectively tests players’ theory-of-mind (ToM) abilities (Kosinski, 2023; Mao et al., 2023). In behavioral economics, the Cognitive Hierarchy Model (Camerer et al., 2004) categorizes players as follows: Level-0 players choose numbers randomly. Level-1 players assume others are Level-0 and pick two-thirds of an expected mean of 50. Level- k players believe that the participants include levels 0 to $k - 1$, and therefore choose $(2/3)^k \times 50$. The optimal outcome is to choose 0 for all players, achieving a Nash equilibrium. In this game, $N = 5$ agents using same backbone model with the same prompting method (e.g., SELFGOAL) play games with each other to observe group behavior.

First-price Auction: AucArena We use **AucArena** (Chen et al., 2023) as the implementation of first-price auctions. An auctioneer collects and announces the bids of all participants, revealing the current highest bid. Participants must publicly make their decisions after privately considering their bids. The auction comprises if $K = 15$ items with values ranging from \$2,000 to \$10,000, with an increment of \$2,000 between each item. These items are presented in a randomized sequence, making the auction last for $K = 15$ rounds. $N = 4$ agents participate in the auction as bidders. Each agent aims to secure the highest profit by the end of the auction and thereby outperform all competitors. In our experiment, we set the budget for each bidder at \$20,000. We have an agent, enhanced by various methods (e.g., SELFGOAL), using different backbone models to compete against three identical opponents powered by the same model (GPT-3.5 (OpenAI, 2022)).

Bargaining: DealOrNotDeal We use **DealOrNotDeal** (Lewis et al., 2017) to im-

plement the bargaining over multiple issues. $N = 2$ agents, namely Alice and Bob, are presented with sets of items (e.g., books, hats, balls) and must negotiate their distribution. Each agent is randomly assigned an integer value between 0 and 10 for each item, ensuring that the total value of all items for any agent does not exceed 10. The bargaining goes on for $K = 10$ rounds, and if the agents fail to agree on the distribution of items within 10 rounds, neither party profits. The goal is to minimize profit discrepancies between the two agents. We randomly select $M = 50$ items for Alice and Bob to negotiate over. The final profits at the end of the negotiation for Alice and Bob are defined as P_{Alice} and P_{Bob} , respectively. Note that, we alter the prompting methods of the agent behind Alice, and keep Bob fixed (GPT-3.5).

4.2 Agent Framework Baselines and Backbone LLMs

We adopt two types of agent frameworks providing guidance for achieving high-level goals in the above tasks.¹ One is **task decomposition** framework, including ReAct (Yao et al., 2023) and ADAPT (Prasad et al., 2024). ReAct enables agents to reason before acting, while ADAPT recursively plans and decomposes complex sub-tasks when the LLM cannot execute them. Another is **experience summarization** framework, including Reflexion (Shinn et al., 2023) and CLIN (Majumder et al., 2023). Reflexion prompts agents to reflect on failed task attempts and retry. CLIN creates a memory of causal abstractions to assist trials in future by reflecting on past experiences, expressed as “A [may/should] be necessary for B.”. To drive these language agent frameworks, we use the following LLMs: **GPT-3.5-Turbo** (gpt-3.5-turbo-1106) (OpenAI, 2024) and **GPT-4-Turbo** (gpt-4-1106-preview) (OpenAI, 2024); **Gemini 1.0 Pro** (Team et al., 2023); **Mistral-7B-Instruct-v0.2** (Jiang et al., 2023) and a Mixture of Experts (MoE) model **Mixtral-8x7B-Instruct-v0.1** (Jiang et al., 2024); **Qwen 1.5** (7B and 72B variants) (Bai et al., 2023). The temperature is set to 0 to minimize randomness.

4.3 Metrics for Tasks

In GAMA-Bench’s Public Goods Game (Huang et al., 2024), where N players participating in re-

¹Implementation details are in Appendix A.4.

peated T times, the score S_1 for this game is then given by: $S_1 = \frac{1}{NT} \sum_{ij} C_{i,j}$, where $C_{i,j} \in [0, 1]$ is the proposed contribution of player i in round j .

In GAMA-Bench’s Guess 2/3 of the Average Game (Huang et al., 2024), the score S_2 is calculated by $S_2 = 100 - \frac{1}{NT} \sum_{ij} C_{i,j}$, where $C_{i,j}$ is the number chosen by player i in round j .

In AucArena’s First-price Auction (Chen et al., 2023), we use the TrueSkill Score (Herbrich et al., 2006; Minka et al., 2018) (Appendix A.5) to rank the profits of agents. TrueSkill Score estimates dynamic skill levels (μ) through Bayesian statistics while considering the uncertainty (σ) in their true skills. Thus the performance score of an agent is defined as $S_3 = \text{TrueSkill Score}$. This method is commonly used in competitions such as online games or tournaments.

In DealOrNotDeal’s Bargaining Game (Lewis et al., 2017), we calculate the absolute difference in their profits: $S_4 = \frac{|P_{Alice} - P_{Bob}|}{M}$, where P_{Alice}, P_{Bob} represents the profits at the end of the negotiation, and M is the number of items to negotiate on. (S_4 can also be represented by TrueSkill Score for convenience.)

5 Results and Analysis

5.1 Main Results

The main results across 4 scenarios are presented in Table 2. Overall, our SELFGOAL significantly outperforms all baseline frameworks in various environments containing high-level goals, where larger LLMs produce higher gains. When diving into the generated guidelines and corresponding agents’ behaviors, we find that some of those subgoals given by task decomposition methods like ReAct and ADAPT are no longer suited for the current situation. For example, “bid on the most expensive item” is not useful when the budget is tight. Moreover, task decomposition before interacting with the environment does not consider the practical experience, leading to broad and meaningless guidance. For example, in Public Goods Game, ADAPT provides broad subgoals like “It’s important to strike a balance between contributing enough tokens to the public pot to earn a significant payoff while retaining enough tokens in my private collection for future rounds”. In contrast, post-hoc experience summarization methods, i.e., Reflexion and CLIN,

tend to induce too detailed guidelines, lacking a correlation with the main goal and might deviating agents from their paths. For example, CLIN produces subgoals focusing on minutiae, such as “Considering the distribution of numbers chosen by opponents may be necessary to make an informed decision on your own selection.”

In comparison, SELFGOAL overcomes both of the shortcomings. At each round, SELFGOAL decomposes new nodes referring to existing guidance, aligning with the main goal as the game progresses. For example, in Public Good Game, the initial subgoal is “The player aims to contribute strategically based on their assessment of other players’ behaviors and the overall distribution of tokens in the public pot.” If all players contribute less to the public pot during the game, SELFGOAL absorbs the observation and refines existing nodes to “If the player notices that the average contribution of the group has been increasing in recent rounds, they might choose to contribute fewer tokens in the current round to avoid over-contributing and potentially losing out on their own gain.” According to the new subgoal as a practical guideline, agents can dynamically adjust their contributions.²

Interestingly, SELFGOAL shows superior performance in smaller LLMs as well, while others can not due to the deficiency of induction and summarization capability of these models. For example, CLIN is 0.7 inferior to Reflexion for Mistral-7B and 5.77 for Qwen-7B in Guess 2/3 of the Average, but SELFGOAL brings improvements consistently. This can be attributed to the logical, structural architecture of GOALTREE in SELFGOAL. At each time for decomposition, the model receives existing subgoals on the last layer of GOALTREE as clear references, making it easy for decomposition.

Competition between Different Agent Framework Previous results are mostly evaluated against a fixed baseline (GPT-3.5). To understand how these agent frameworks behave when competing with each other, we set an AucArena for

²More details of GOALTREE are in Appendix A.7.

Table 2: Comparison of the SELFGOAL powered by different models with alternative methods across four scenarios. The best results are **bolded**, and the second best ones are underlined.

Methods	ReAct	ADAPT	Reflexion	CLIN	SELFGOAL	ReAct	ADAPT	Reflexion	CLIN	SELFGOAL
Public Goods Game: GAMA (Huang et al., 2024) ($S_1 \downarrow$)						Guess 2/3 of the Average: GAMA (Huang et al., 2024) ($S_2 \uparrow$)				
Mistral-7B	55.70	46.00	51.28	41.00	28.45	89.43	84.91	<u>92.65</u>	91.95	93.64
Mixtral-8x7B	46.05	55.80	34.65	52.69	32.00	82.16	79.46	89.73	74.33	89.50
Qwen-7B	66.55	56.44	60.15	<u>55.59</u>	54.93	65.11	55.95	<u>69.99</u>	64.22	72.99
Qwen-72B	<u>20.75</u>	22.95	21.57	24.60	8.45	78.87	88.77	<u>91.47</u>	83.65	94.51
Gemini Pro	37.55	25.78	34.00	39.20	19.20	77.90	73.45	71.82	76.58	<u>77.33</u>
GPT-3.5	61.20	<u>42.25</u>	46.95	47.15	42.19	73.44	64.14	<u>78.75</u>	63.25	83.28
GPT-4	19.55	<u>16.70</u>	22.90	31.35	11.95	92.57	91.31	<u>94.41</u>	90.88	94.54
Methods	ReAct	ADAPT	Reflexion	CLIN	SELFGOAL	ReAct	ADAPT	Reflexion	CLIN	SELFGOAL
First-price Auction: AucArena (Chen et al., 2023) ($S_3 \uparrow$)						Bargaining: DealOrNotDeal (Lewis et al., 2017) ($S_4 \downarrow$)				
Mistral-7B	23.91	23.03	<u>26.24</u>	24.27	28.21	2.57	2.38	<u>1.97</u>	2.32	1.88
Mixtral-8x7B	35.85	32.35	33.18	<u>36.37</u>	39.23	2.38	2.66	2.46	<u>2.34</u>	1.97
Qwen-7B	29.88	30.15	32.97	<u>33.44</u>	33.50	2.83	2.88	3.15	<u>2.73</u>	2.05
Qwen-72B	34.77	34.25	<u>35.92</u>	<u>34.24</u>	36.48	2.59	2.10	<u>2.06</u>	<u>2.26</u>	2.00
Gemini Pro	36.12	36.47	38.82	36.79	39.28	2.10	2.33	2.28	2.36	1.95
GPT-3.5	<u>22.85</u>	22.10	22.00	21.21	27.40	<u>2.31</u>	2.95	2.44	2.87	2.20
GPT-4	36.46	35.40	34.41	<u>38.98</u>	39.02	1.94	<u>1.80</u>	1.92	1.83	1.71

Table 3: Result of auction competitions between the reported five agents with baseline frameworks and our SELFGOAL.

Methods	ReAct	ADAPT	Reflexion	CLIN	SELFGOAL
GPT-3.5	23.96 $\pm$ 1.72	20.46 $\pm$ 1.79	<u>25.72$\pm$1.71</u>	22.95 $\pm$ 1.73	29.59$\pm$1.99
GPT-4	22.62 $\pm$ 1.80	24.85 $\pm$ 1.78	21.79 $\pm$ 1.79	<u>27.16$\pm$1.74</u>	28.98$\pm$1.88

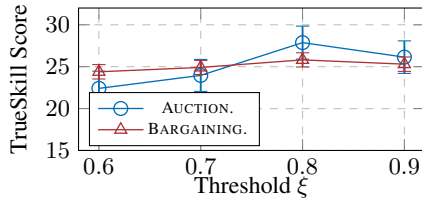


Figure 2: Granularity control of the threshold ξ in SELFGOAL’s stopping mechanism.

a multi-agent comparison. As shown in Table 3, SELFGOAL has a clear advantage over baselines. When looking closer at the bidding behaviors, we find that other methods tend to be overly cautious. They often stop bidding or avoid participating once bidding starts, resulting in zero profits. However, SELFGOAL takes a different approach by bidding frequently in the early stages of the bidding war when competition is less intense. This allows for purchasing high-priority items early on, avoiding fierce competition in the later stages.

5.2 Analysis of SELFGOAL

How does the granularity of guidelines in GOAL-TREE affect task solving? As discussed in §5.1, SELFGOAL adjusts to the dynamic environment by setting different depths, where subgoal nodes of deeper layers provide more detailed instructions.

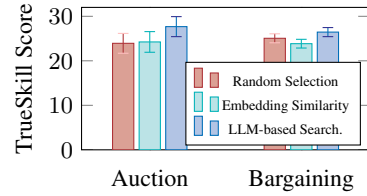


Figure 3: Ablation study of different search modules.

Here, we explore how such granularity affects the performance of SELFGOAL. We use Auction and Bargaining environments as testbeds, and modify the level of subgoals by setting the threshold ξ in the stopping mechanism as 0.6, 0.7, 0.8, and 0.9. According to Figure 2, the agent’s performance initially improves with increasing depth but eventually diminishes. A shallow tree ($\xi = 0.6$) lacks guidance details, thus leading to the poorest performance. Yet, the deepest tree ($\xi = 0.9$) does not show superior performance, probably because repetitive guidance interferes with model selection of useful guidance. Redundant nodes increase the candidate set, making it difficult for the search module to select all the valuable nodes. In fact, the search module always focuses on multiple nodes representing the same meaning, resulting in the loss of other helpful nodes. This experiment confirms that more detailed instructions help language agents achieve high-level goals, but only with a balanced, adaptive depth of the guidance tree to mitigate the drawbacks of overly detailed guidance (We further conduct a case study (Appendix A.1) to demonstrate how SELFGOAL’s focus on granularity control provides distinct advantages).

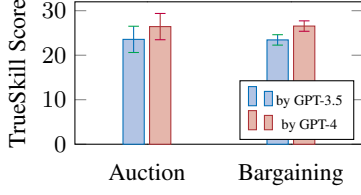


Figure 4: Ablation study of the model that generates GOALTREE, either by a stronger (GPT-4) or weaker (GPT-3.5) model. The rest of the agent framework is driven by GPT-3.5.

Ablation Study of Search Module *Can the Search Module in SELFGOAL succeed in finding useful subgoal nodes?* We employ two methods as baselines to replace the original LLM-based search module, which is instantiated with GPT-3.5. One baseline is *random selection*, where we randomly choose an node from the set of subgoal nodes. The other is the selection based on *embedding similarity*, which selects the subgoals most similar to the current situation based on cosine similarity. On multi-round games as Auction and Bargaining, we keep the Trueskill Score for evaluating the rankings of these methods. As shown in Figure 3, the LLM search module gains a better score in both games. Besides, similarity-based method performs worse than random selection in Bargaining, which could be the reason that the guidance is usually short, making it hard to capture semantic embeddings between subgoals and situations. This experiment demonstrates the rationality of the LLM-based search module in SELFGOAL’s design.

How does the quality of GOALTREE affect goal achievement? To explore the influence of GOALTREE on SELFGOAL, we conduct an experiment in Auction and Bargaining Games by replacing the model that constructs GOALTREE with GPT-4 or GPT-3.5 for comparison, while keeping the model that utilizes the tree fixed as GPT-3.5. Results in Figure 4 illustrate that higher-quality GOALTREE (from GPT-4) significantly boosts the performance of SELFGOAL, with gains of +2.87 in Auction and +3.10 in Bargaining compared to one using GPT-3.5. This improvement comes from more abundant and higher-quality guidance, generated by a strong model equipped with better understanding and summarizing capabilities (We also conduct an ablation study on the impact of pruning on GOALTREE in Appendix A.2).

Can SELFGOAL improve the rationality in agents’ behaviors? Aside from the final perfor-

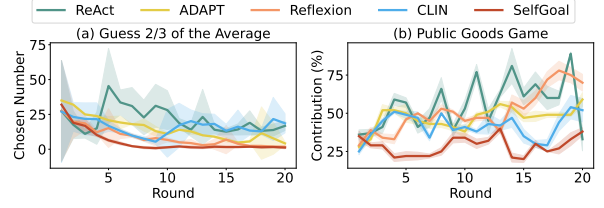


Figure 5: Patterns of model behavior in repeated games. (a): Adjustments in number predictions within the Guessing Game. Our SELFGOAL shows improved ToM abilities by converging to a guess of zero more quickly in each round. (b): Fluctuations in contributions within the Public Goods game. The agent equipped with SELFGOAL displays more rational behavior (*i.e.*, achieving a Nash equilibrium) by consistently contributing fewer tokens than other methods.

mance gain, we are also interested in whether each agent behavior at every turn benefits from SELFGOAL. Therefore, we use two games from GAMA-Bench to examine the impact of SELFGOAL on model behavior, where behavioral changes are easier to evaluate. Here, we use LLMs with great improvement from SELFGOAL, *i.e.*, Mistral-7B for Public Goods Game and Qwen-72B for Guessing 2/3 Average Number Game. We record patterns in the model’s number predictions and token contributions by visualizing data from 20 repeated experiments. Note that GOALTREE is updated across these 20 rounds of games. With SELFGOAL, agents in the Public Goods scenario consistently act more rationally compared to those using alternative methods, as illustrated in Figure 5(a). For the Guessing Game, enhanced models showed smoother, steadily declining curves, indicating faster convergence to the Nash equilibrium (Figure 5(b)).

6 Conclusion

In this paper, we introduce SELFGOAL, an agent framework that enhances the capabilities of LLMs for achieving high-level goals across various dynamic tasks and environments. We demonstrate that SELFGOAL significantly improves agent performance by dynamically generating and refining a hierarchical GOALTREE of contextual subgoals based on interactions with the environments. Experiments show that this method is effective in both competitive and cooperative scenarios, outperforming baseline approaches. Moreover, GOALTREE can be continually updated as agents with SELFGOAL further engage with the environments, enabling them to navigate complex environments with greater precision and adaptability.

Limitation

SELF GOAL does incur higher computational costs compared to the baseline methods within a reasonable range. Specifically, SELF GOAL requires approximately five times the computational resources of the baseline methods, as shown in Table 5. However, this investment produces a significant performance boost, as SELF GOAL achieves a TrueSkill improvement of +5.9 over ReAct. In contrast, hierarchical methods like ADAPT require four times the baseline computational resources but do not enhance performance. This clear difference underscores the efficiency of our approach, demonstrating that the extra resources are effectively utilized to produce meaningful improvements.

Furthermore, while SELF GOAL is effective for smaller models, we recognize that its performance may be constrained by the inherent limitations of models in understanding and summarizing the capabilities, which could prevent SELF GOAL from reaching its full potential.

References

- Michael Ahn, Anthony Brohan, Noah Brown, Yevgen Chebotar, Omar Cortes, Byron David, Chelsea Finn, Chuyuan Fu, Keerthana Gopalakrishnan, Karol Hausman, Alex Herzog, Daniel Ho, Jasmine Hsu, Julian Ibarz, Brian Ichter, Alex Irpan, Eric Jang, Rosario Jauregui Ruano, Kyle Jeffrey, Sally Jesmonth, Nikhil J Joshi, Ryan Julian, Dmitry Kalashnikov, Yuheng Kuang, Kuang-Huei Lee, Sergey Levine, Yao Lu, Linda Luu, Carolina Parada, Peter Pastor, Jornell Quiambao, Kanishka Rao, Jarek Rettinghouse, Diego Reyes, Pierre Sermanet, Nicolas Sievers, Clayton Tan, Alexander Toshev, Vincent Vanhoucke, Fei Xia, Ted Xiao, Peng Xu, Sichun Xu, Mengyuan Yan, and Andy Zeng. 2022. [Do as i can, not as i say: Grounding language in robotic affordances](#).
- Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, Binyuan Hui, Luo Ji, Mei Li, Junyang Lin, Runji Lin, Dayiheng Liu, Gao Liu, Chengqiang Lu, Keming Lu, Jianxin Ma, Rui Men, Xingzhang Ren, Xuancheng Ren, Chuanqi Tan, Sinan Tan, Jianhong Tu, Peng Wang, Shijie Wang, Wei Wang, Shengguang Wu, Benfeng Xu, Jin Xu, An Yang, Hao Yang, Jian Yang, Shusheng Yang, Yang Yao, Bowen Yu, Hongyi Yuan, Zheng Yuan, Jianwei Zhang, Xingxuan Zhang, Yichang Zhang, Zhenru Zhang, Chang Zhou, Jingren Zhou, Xiaohuan Zhou, and Tianhang Zhu. 2023. [Qwen technical report](#).
- Andrew G Barto and Sridhar Mahadevan. 2003. Recent advances in hierarchical reinforcement learning. *Discrete event dynamic systems*, 13:341–379.
- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*.
- Colin F. Camerer, Ho Teck-Hua, and Chong Juin-Kuan. 2004. A cognitive hierarchy model of games. *The Quarterly Journal of Economics*, 119.
- Jiangjie Chen, Siyu Yuan, Rong Ye, Bodhisattwa Prasad Majumder, and Kyle Richardson. 2023. [Put your money where your mouth is: Evaluating strategic planning and execution of llm agents in an auction arena](#).
- Constantinos Daskalakis, Paul W Goldberg, and Christos H Papadimitriou. 2009. The complexity of computing a nash equilibrium. *Communications of the ACM*, 52(2):89–97.
- Kutluhan Erol, James Hendler, and Dana S Nau. 1994. Htn planning: Complexity and expressivity. In *AAAI*, volume 94, pages 1123–1128.
- Valerie Goffaux, Judith Peters, Julie Haubrechts, Christine Schiltz, Bernadette Jansma, and Rainer Goebel. 2011. [From coarse to fine? spatial and temporal dynamics of cortical face processing](#). *Cerebral Cortex*, page 467–476.
- Ralf Herbrich, Tom Minka, and Thore Graepel. 2006. Trueskill™: a bayesian skill rating system. *Advances in neural information processing systems*, 19.
- Christopher Hoang, Sungryull Sohn, Jongwook Choi, Wilka Carvalho, and Honglak Lee. 2021. [Successor feature landmarks for long-horizon goal-conditioned reinforcement learning](#). In *Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual*, pages 26963–26975.
- Jen-tse Huang, Eric John Li, Man Ho Lam, Tian Liang, Wenxuan Wang, Youliang Yuan, Wenxiang Jiao, Xing Wang, Zhaopeng Tu, and Michael R Lyu. 2024. [How far are we on the decision-making of llms? evaluating llms’ gaming ability in multi-agent environments](#). *ArXiv preprint*, abs/2403.11807.
- Wenlong Huang, Pieter Abbeel, Deepak Pathak, and Igor Mordatch. 2022a. [Language models as zero-shot planners: Extracting actionable knowledge for embodied agents](#). In *International Conference on*

744	<i>Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA</i> , volume 162 of <i>Proceedings of Machine Learning Research</i> , pages 9118–9147. PMLR.	
745		
746		
747		
748	Wenlong Huang, Fei Xia, Ted Xiao, Harris Chan, Jacky Liang, Pete Florence, Andy Zeng, Jonathan Tompson, Igor Mordatch, Yevgen Chebotar, Pierre Sermanet, Noah Brown, Tomas Jackson, Linda Luu, Sergey Levine, Karol Hausman, and Brian Ichter. 2022b. Inner monologue: Embodied reasoning through planning with language models .	
749		
750		
751		
752		
753		
754		
755	Zhiao Huang, Fangchen Liu, and Hao Su. 2019. Mapping state space using landmarks for universal goal reaching . In <i>Advances in Neural Information Processing Systems 32: Annual Conference on Neural Information Processing Systems 2019, NeurIPS 2019, December 8-14, 2019, Vancouver, BC, Canada</i> , pages 1940–1950.	
756		
757		
758		
759		
760		
761		
762	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, L��lio Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timoth��e Lacroix, and William El Sayed. 2023. Mistral 7b .	
763		
764		
765		
766		
767		
768		
769	Albert Q Jiang, Alexandre Sablayrolles, Antoine Roux, Arthur Mensch, Blanche Savary, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Emma Bou Hanna, Florian Bressand, et al. 2024. Mixtral of experts . <i>ArXiv preprint</i> , abs/2401.04088.	
770		
771		
772		
773		
774	Tushar Khot, Harsh Trivedi, Matthew Finlayson, Yao Fu, Kyle Richardson, Peter Clark, and Ashish Sabharwal. 2022. Decomposed prompting: A modular approach for solving complex tasks . <i>ArXiv preprint</i> , abs/2210.02406.	
775		
776		
777		
778		
779	Tushar Khot, Harsh Trivedi, Matthew Finlayson, Yao Fu, Kyle Richardson, Peter Clark, and Ashish Sabharwal. 2023. Decomposed prompting: A modular approach for solving complex tasks .	
780		
781		
782		
783	Michal Kosinski. 2023. Theory of mind might have spontaneously emerged in large language models . <i>ArXiv preprint</i> , abs/2302.02083.	
784		
785		
786	Alain Ledoux. 1981. Concours r��sultats complets: Les victimes se sont plu �� jouer le 14 d’atout. <i>Jeux & Strat��gie</i> , 2(10):10–11.	
787		
788		
789	Mike Lewis, Denis Yarats, Yann Dauphin, Devi Parikh, and Dhruv Batra. 2017. Deal or no deal? end-to-end learning of negotiation dialogues . In <i>Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing</i> , pages 2443–2453, Copenhagen, Denmark. Association for Computational Linguistics.	
790		
791		
792		
793		
794		
795		
796	Bill Yuchen Lin, Yicheng Fu, Karina Yang, Faeze Brahman, Shiyu Huang, Chandra Bhagavatula, Prithviraj Ammanabrolu, Yejin Choi, and Xiang Ren. 2023. Swiftsage: A generative agent with fast and slow thinking for complex interactive tasks .	
797		
798		
799		
800		
	Bo Liu, Yuqian Jiang, Xiaohan Zhang, Qiang Liu, Shiqi Zhang, Joydeep Biswas, and Peter Stone. 2023. Llm+p: Empowering large language models with optimal planning proficiency .	801
		802
		803
		804
	Yuchen Liu, Luigi Palmieri, Sebastian Koch, Ilche Georgievski, and Marco Aiello. 2024. Delta: Decomposed efficient long-term robot task planning using large language models .	805
		806
		807
		808
	Chengdong Ma, Ziran Yang, Minquan Gao, Hai Ci, Jun Gao, Xuehai Pan, and Yaodong Yang. 2024. Red teaming game: A game-theoretic framework for red teaming language models .	809
		810
		811
		812
	Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegreffe, Uri Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, Shashank Gupta, Bodhisattwa Prasad Majumder, Katherine Hermann, Sean Welleck, Amir Yazdanbakhsh, and Peter Clark. 2023. Self-refine: Iterative refinement with self-feedback .	813
		814
		815
		816
		817
		818
		819
	Bodhisattwa Prasad Majumder, Bhavana Dalvi Mishra, Peter Jansen, Oyvind Tafjord, Niket Tandon, Li Zhang, Chris Callison-Burch, and Peter Clark. 2023. Clin: A continually learning language agent for rapid task adaptation and generalization .	820
		821
		822
		823
		824
	Yuanyuan Mao, Shuang Liu, Pengshuai Zhao, Qin Ni, Xin Lin, and Liang He. 2023. A review on machine theory of mind .	825
		826
		827
	Tom Minka, Ryan Clevon, and Yordan Zaykov. 2018. Trueskill 2: An improved bayesian skill rating system . <i>Technical Report</i> .	828
		829
		830
	Kolby Nottingham, Bodhisattwa Prasad Majumder, Bhavana Dalvi Mishra, Sameer Singh, Peter Clark, and Roy Fox. 2024. Skill set optimization: Reinforcing language model behavior via transferable skills .	831
		832
		833
		834
	OpenAI. 2022. Chatgpt .	835
	OpenAI. 2024. Gpt-4 technical report .	836
	Debjit Paul, Mete Ismayilzada, Maxime Peyrard, Beatriz Borges, Antoine Bosselut, Robert West, and Boi Faltings. 2024. Refiner: Reasoning feedback on intermediate representations .	837
		838
		839
		840
	Damien Pellier, Alexandre Albore, Humbert Fiorino, and Rafael Bailon-Ruiz. 2023. Hddl 2.1: Towards defining a formalism and a semantics for temporal htn planning .	841
		842
		843
		844
	Archiki Prasad, Alexander Koller, Mareike Hartmann, Peter Clark, Ashish Sabharwal, Mohit Bansal, and Tushar Khot. 2024. Adapt: As-needed decomposition and planning with language models .	845
		846
		847
		848
	Faisal Rahutomo, Teruaki Kitasuka, Masayoshi Aritsugi, et al. 2012. Semantic cosine similarity. In <i>The 7th international student conference on advanced science and technology ICAST</i> , volume 4, page 1. University of Seoul South Korea.	849
		850
		851
		852
		853

854	Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. 2023. Reflexion: Language agents with verbal reinforcement learning .	910
855		911
856		912
857		913
858	Ishika Singh, David Traum, and Jesse Thomason. 2024. Twostep: Multi-agent task planning using classical planners and large language models . <i>ArXiv preprint</i> , abs/2403.17246.	914
859		915
860		916
861		917
862	Gemini Team, Rohan Anil, Sebastian Borgeaud, Yonghui Wu, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, et al. 2023. Gemini: a family of highly capable multimodal models . <i>ArXiv preprint</i> , abs/2312.11805.	918
863		
864		
865		
866		
867		
868	Guanzhi Wang, Yuqi Xie, Yunfan Jiang, Ajay Mandlekar, Chaowei Xiao, Yuke Zhu, Linxi (Jim) Fan, and Anima Anandkumar. 2023a. Voyager: An open-ended embodied agent with large language models . <i>ArXiv preprint</i> , abs/2305.16291.	919
869		920
870		921
871		922
872		
873	Lei Wang, Wanyu Xu, Yihuai Lan, Zhiqiang Hu, Yunshi Lan, Roy Ka-Wei Lee, and Ee-Peng Lim. 2023b. Plan-and-solve prompting: Improving zero-shot chain-of-thought reasoning by large language models . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 2609–2634, Toronto, Canada. Association for Computational Linguistics.	923
874		924
875		925
876		926
877		927
878		
879		
880		
881	Ruoyao Wang, Peter Jansen, Marc-Alexandre Côté, and Prithviraj Ammanabrolu. 2022. ScienceWorld: Is your agent smarter than a 5th grader? In <i>Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing</i> , pages 11279–11298, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.	
882		
883		
884		
885		
886		
887		
888	Jian Xie, Kai Zhang, Jiangjie Chen, Tinghui Zhu, Renze Lou, Yuandong Tian, Yanghua Xiao, and Yu Su. 2024. Travelplanner: A benchmark for real-world planning with language agents . <i>ArXiv preprint</i> , abs/2402.01622.	
889		
890		
891		
892		
893	Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. 2023. React: Synergizing reasoning and acting in language models .	
894		
895		
896		
897	Siyu Yuan, Jiangjie Chen, Ziquan Fu, Xuyang Ge, Soham Shah, Charles Jankowski, Yanghua Xiao, and Deqing Yang. 2023. Distilling script knowledge from large language models for constrained language planning . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 4303–4325, Toronto, Canada. Association for Computational Linguistics.	
898		
899		
900		
901		
902		
903		
904		
905	Yikai Zhang, Siyu Yuan, Caiyu Hu, Kyle Richardson, Yanghua Xiao, and Jiangjie Chen. 2024. Timearena: Shaping efficient multitasking language agents in a time-aware simulation . <i>ArXiv preprint</i> , abs/2402.05733.	
906		
907		
908		
909		
	Andrew Zhao, Daniel Huang, Quentin Xu, Matthieu Lin, Yong-Jin Liu, and Gao Huang. 2024. Expel: Llm agents are experiential learners . In <i>Thirty-Eighth AAAI Conference on Artificial Intelligence, AAAI 2024, Thirty-Sixth Conference on Innovative Applications of Artificial Intelligence, IAAI 2024, Fourteenth Symposium on Educational Advances in Artificial Intelligence, EAAI 2024, February 20-27, 2024, Vancouver, Canada</i> , pages 19632–19642. AAAI Press.	
	Yi Zheng, Chongyang Ma, Kanle Shi, and Haibin Huang. 2023. Agents meet okr: An object and key results driven agent system with hierarchical self-collaboration and self-evaluation .	
	Xuhui Zhou, Hao Zhu, Leena Mathur, Ruohong Zhang, Haoifei Yu, Zhengyang Qi, Louis-Philippe Morency, Yonatan Bisk, Daniel Fried, Graham Neubig, and Maarten Sap. 2024. Sotopia: Interactive evaluation for social intelligence in language agents .	

A SELFGOAL Details

A.1 Case Study

To illustrate how agents from different frameworks reason and plan in a dynamic environment, we conduct a case study using Mistral-7B, a small LLM, as the backbone in a bargaining game (Figure 6). We find that SELFGOAL’s emphasis on granularity control offers clear advantages. SELFGOAL provides agents with actionable guidance such as “ask clarifying questions”, prompting agents to pay early attention to their opponent’s psychological assessment and different valuations of items. After acquiring a partner’s valuation, SELFGOAL then gives guidance such as “make concessions”, leading the agent to propose a plan that gives up a particular item in exchange for minimizing the profit difference.

In contrast, CLIN advises agents to “consider the preference of the partner”, which leads agents to focus on the opponent’s preferences, but may result in plans that sacrifice their own interests to improve the other party’s income. ADAPT, which decomposes tasks beforehand, provides very broad advice such as “equal allocation”. This generic advice aims to minimize the profit gap but may not be suitable for scenarios lacking knowledge of the partner’s valuation. Consequently, the model proposes allocation plans without first clarifying the partner’s valuations, assuming that all participants have the same valuation for each item.

A.2 Does pruning the GOALTREE affect search quality?

Table 4: Comparison of agents guided by GOALTREE with and without pruning.

GOALTREE	Scenario	
	Auction	Bargaining
Pruned	24.74 $\pm$ 3.22	24.90 $\pm$ 1.21
w/o Pruned	25.25 $\pm$ 3.23	25.09 $\pm$ 1.21

We investigate whether pruning nodes not selected for a long time from the target tree affects the Search Module’s decisions. Pruning begins after the Decompose Module completes building the tree, and nodes unselected for more than five consecutive rounds will be deleted. We assess the impact of pruning on GPT-3.5’s performance in Auction and Bargaining. As shown in Table 4,

the TrueSkill Score with and without pruning are similar. This suggests that nodes not chosen for extended periods do not compromise the Search Module’s decision-making effectiveness. This efficiency likely results from our Search Module using prior knowledge from LLM to identify and avoid selecting unnecessary nodes, akin to lazy deletion. For efficiency, these redundant nodes are also removed every five rounds.

A.3 Computational Efficiency Analysis

Table 5: Computational Efficiency of Different Methods in Auction Per Round.

Method	OpenAI Cost	Tokens Used	Computation Time	Performance
ReAct	0.366	295,556.6	5.42 min	22.90
ADAPT	1.248	834,382.7	8.28 min	22.30
Reflexion	0.434	359,674.8	5.41 min	22.32
CLIN	0.448	372,803.4	5.52 min	21.41
SELFGOAL	2.20	1,651,328.8	13.46 min	28.81

We evaluated the computational efficiency of SELFGOAL by conducting experiments in the Auction Arena over 5 rounds, using GPT-3.5 as the backbone model. We monitored the average OpenAI cost, tokens used, and computation time per round. As shown in Table 5, although SELFGOAL incurred higher costs and computation times, these were within an acceptable range and significantly improved model performance, as evidenced by the TrueSkill metric.

A.4 Implementation Details

We compare our SELFGOAL with the following methods: ReAct (Yao et al., 2023), which induces an LLM actor to engage in preliminary reasoning about the task before initiating action, Reflexion (Shinn et al., 2023), which encourages an LLM actor to re-assess unsuccessful task attempts before attempting the task again, CLIN (Majumder et al., 2023), which leverages historical insights to deduce transition strategies, articulated as “A [may/should] be necessary for A”. To adapt these methods to our experimental environment, we update the memory of the CLIN/Reflexion approach at each timestep within a single trial, whether it is a bid in the Auction environment, a dialogue round in the Negotiation environment, or a game round in GAMA-Bench. Specifically, for Reflexion, the model uses historical steps from the current trial to generate verbal self-reflections. These self-reflections are then added to long-term memory, providing valuable feedback for future trials. In the case of CLIN, we use the BASE method due to the absence of a training set in our environment. The memory is updated at each step by prompting the model

Task: Bargaining (DealOrNotDeal)		Item Valuation			
Task: You are Alice and you want to negotiate with your partner, Bob, about the allocation plan of items: 1 book, 1 hat, and 2 ball. Your goal is to minimize the profit gap. Observation: Bob: Hi, Let's start to negotiate about the allocation plan.			Book	Hat	Ball
		Alice	0	6	2
		Bob	8	2	0

CLIN	SelfGoal
Learnings: Understanding the other party's primary objectives SHOULD BE NECESSARY to propose mutually beneficial agreements. Action: Let's start by discussing our priorities and preferences for the items. Observation: Bob: I prefer to receive the book and the ball. Action: Great! I will take 2 balls and you can have the book and ball. Profit: Alice: \$ 4, Bob \$ 10	Sub-goal: Ask clarifying questions: [...] If necessary, ask Bob to clarify or provide more details about his valuation of each item. [...] Action: [...] Based on my understanding, you have a different valuation for each item compared to me. Could you please share your valuation for each item? [...] Observation: Bob: I value book at \$8, hat at \$2, book at \$0. Sub-goal: Be willing to compromise: Be willing to make concessions and adjust your proposals to accommodate your partner's feedback and suggestions. Action: Great! Since I value book at \$0, hat at \$6, ball at \$2, to minize our profit gap, I propose I take 1 hat and 1 ball while you take the book. Then we both give up 1 ball. Observation: Bob: Great! I agree to that proposal. Profit: Alice: \$ 8, Bob \$ 8

ADAPT
PLAN: I will propose an allocation plan that give an Equal allocation of items. [...] Action: [...] I'd like to propose that we split the items evenly, with you taking 1 book, and 1 balls, and me taking 1 hat, and 1 ball. [...] Observation: Bob: Great! I agree to that proposal. Profit: Alice: \$ 8, Bob \$ 10

Figure 6: In the Bargaining task, Mistral-7B with CLIN or ADAPT gives guidance that is either too broad or too detailed resulting in large profit discrepancy, whereas SELFgoal is successful.

with historical steps from the current trial and all previous memories to generate an updated memory, which includes a new list of semi-structured causal abstractions. This updated memory is then incorporated into the historical memories.

A.5 Details of TrueSkill Score

In a game with a population of n players $\{1, \dots, n\}$, consider a match where k teams compete. The team assignments are specified by k non-overlapping subsets $A_j \subset \{1, \dots, n\}$ of the player population, with $A_i \cap A_j = \emptyset$ for $i \neq j$. The outcome $\mathbf{r} := (r_1, \dots, r_k) \in \{1, \dots, k\}$ is defined by a rank r_j for each team j , with $r = 1$ indicating the winner and draws possible when $r_i = r_j$. Ranks are based on the game's scoring rules.

The probability $P(\mathbf{r} \mid \mathbf{s}, A)$ of the game outcome $\mathbf{r}$ is modeled given the skills $\mathbf{s}$ of the participating players and the team assignments $A := \{A_1, \dots, A_k\}$. From Bayes' rule, we get the posterior distribution

$$p(\mathbf{s} \mid \mathbf{r}, A) = \frac{P(\mathbf{r} \mid \mathbf{s}, A)p(\mathbf{s})}{P(\mathbf{r} \mid A)}.$$

We assume a factorizing Gaussian prior distribution, $p(\mathbf{s}) := \prod_{i=1}^n N(s_i; \mu_i, \sigma_i^2)$. Each player i is assumed to exhibit a performance $p_i \sim N(p_i; s_i, \beta^2)$ in the game, centered around their skill s_i with fixed variance β^2 .

The performance t_j of team j is modeled as the sum of the performances of its members, $t_j := \sum_{i \in A_j} p_i$. Teams are reordered in ascending order

of rank, $r_{(1)} \leq r_{(2)} \leq \dots \leq r_{(k)}$. Disregarding draws, the probability of a game outcome $\mathbf{r}$ is modeled as

$$P(\mathbf{r} \mid \{t_1, \dots, t_k\}) = P(t_{r_{(1)}} > t_{r_{(2)}} > \dots > t_{r_{(k)}})$$

In other words, the order of performances determines the game outcome. If draws are allowed, the winning outcome $r_{(j)} < r_{(j+1)}$ requires $t_{r_{(j)}} > t_{r_{(j+1)}} + \varepsilon$ and the draw outcome $r_{(j)} = r_{(j+1)}$ requires $|t_{r_{(j)}} - t_{r_{(j+1)}}| \leq \varepsilon$, where $\varepsilon > 0$ is a draw margin calculated from the assumed probability of a draw.¹

To report skill estimates after each game, we use an online learning scheme called Gaussian density filtering. The posterior distribution is approximated to be Gaussian and is used as the prior distribution for the next game. If skills are expected to change over time, a Gaussian dynamics factor $N(s_{i,t+1}; s_{i,t}, \gamma^2)$ can be introduced, leading to an additive variance component of γ^2 in the subsequent prior.

Consider a game with $k = 3$ teams with team assignments $A_1 = \{1\}$, $A_2 = \{2, 3\}$ and $A_3 = \{4\}$. Assume that team 1 wins and teams 2 and 3 draw, i.e., $\mathbf{r} := (1, 2, 2)$. The function represented by a factor graph in our case, the joint distribution $p(\mathbf{s}, \mathbf{p}, \mathbf{t} \mid \mathbf{r}, A)$, is given by the product of all the potential functions associated with each factor. The structure of the factor graph provides information about the dependencies of the factors involved and serves as the foundation for efficient inference algorithms. Referring back to Bayes' rule,

the quantities of interest are the posterior distribution $p(s_i | \mathbf{r}, A)$ over skills given game outcome $\mathbf{r}$ and team assignments A . The $p(s_i | \mathbf{r}, A)$ are calculated from the joint distribution by integrating out the individual performances $\{p_i\}$ and the team performances $\{t_i\}$:

$$p(\mathbf{s} | \mathbf{r}, A) = \int_{-\infty}^{\infty} \cdots \int_{-\infty}^{\infty} p(\mathbf{s}, \mathbf{p}, \mathbf{t} | \mathbf{r}, A) d\mathbf{p} d\mathbf{t}.$$

A.6 Instruction Prompt Examples

The instruction prompts of three modules in SELF-GOAL are presented in Listing 1.

Listing 1: The instruction prompts in SELF-GOAL.

Decomposition Instruction:

```
# Main Goal
Humans exhibit numerous behaviors and
sub-goals, which can be traced back to
the primary aim of survival. For
instance:
1. Food Acquisition: To maintain
physical and mental functionality,
individuals seek nourishment. They
target foods with high energy and
nutritional values to augment their
health, thus enhancing survival
possibilities.
2. Shelter Construction: Safe and secure
housing is a fundamental human need. It
offers protection from potentially
harmful natural elements and potential
threats.
```

Imagine you are an agent in a {scene}.

Taking analogy from human behaviors, if your fundamental objective in this scenario is "{goal}", what sub-goals you might have?

```
# Sub-Goal
Here's the current scenario:
```

```
{scene}
```

For the goal: "{sub_goal}", can you further run some deduction for fine-grained goals or brief guidelines?

Search Instruction:

Here's the current scenario:

```
{scene}
```

To better reach your main goal: {objective}, in this context, please do the following:

1. Evaluate how the sub-goals listed below can assist you in reaching your main goal given the present circumstances.

Sub-goals:

```
{guidance}
```

2. Select {width} most useful sub-goals that will help you reach your main goal in the current situation, and note their IDs.

Start by explaining your step-by-step thought process. Then, list the {width} IDs you've chosen, using the format of this example: `{{"IDs": [1, 3, 10, 21, 7]}}`.

Task Solving Instruction:

Here is the current scenarios:

```
{scene}
```

Here are some possible subgoals and guidance derived from your primary objective {main_goal}:

```
{sub_goals}
```

In this round, You may target some of these subgoals and detailed guidance to improve your strategy and action, to achieve your primary objective.

We implemented CLIN and Reflexion methods in our environments as presented in Listing 2.

Listing 2: The instructions for Reflexion and CLIN.

REFLEXION Instruction:

You are an advanced reasoning agent that can improve based on self reflection. Review and reflect on the historical data.

```
{data_log}
```

Based on the history record, in a few sentences, diagnose a possible reason for failure or phrasing discrepancy and devise a new, concise, high level plan that aims to mitigate the same failure. Use complete sentences.

CLIN Instruction:

Review and reflect on the historical data.

```
{data_log}
```

Here are your past learnings:

```
{past_learnings}
```

Based on the history record, formulate or update your learning points that could be advantageous to your strategies

in the future. Your learnings should be strategic, and of universal relevance and practical use for future auctions. Consolidate your learnings into a concise numbered list of sentences. Each numbered item in the list can ONLY be of the form:
X MAY BE NECESSARY to Y.
X SHOULD BE NECESSARY to Y.
X MAY BE CONTRIBUTE to Y.
X DOES NOT CONTRIBUTE to Y.

A.7 Examples of GoalTree

Here, we provide examples of GOALTREE from four environments in Listing 3, with their main goals as follows:

- **Public Goods:** maximize your total token count by the end of the game;
- **Guess 2/3 of the Average:** choose a number that you believe will be closest to 2/3 of the average of all numbers chosen by players, including your selection;
- **First-price Auction:** secure the highest profit at the end of this auction, compared to all other bidders;
- **Bargaining:** minimize the profit gap between yourself and your partner in this negotiation, regardless of your own profit.

Listing 3: Examples of GOALTREE in SELFgoal.

Public Goods Game:
root: Maximize your total token count by the end of the game.
root-0: Maximizing Contribution
root-0-0: Assess the Current State
root-0-0-2: Long-term Token Accumulation
root-0-0-2-3: Collaboration and Competition
root-0-0-2-3-0: Observation and Analysis
root-0-0-2-3-0-1: Identify Potential Collaborators
root-0-0-2-3-0-1-1: Observe Consistency
root-0-0-2-3-0-1-1-1: Establish Trustworthy Partnerships
root-0-0-2-3-0-1-1-1-2: Monitor Trustworthiness
root-0-0-2-3-0-1-1-1-2-1: Identify Unreliable Contributors
root-0-0-2-3-0-1-1-1-2-1-0: Track and Analyze Contributions
root-0-0-2-3-0-1-1-1-2-1-0-1: Identify Inconsistent Contributors
root-0-0-2-3-0-1-1-1-2-1-0-1-1: Monitor Reliability
root-0-0-2-3-0-1-1-1-2-1-0-1-2: Consider Communication
root-0-0-2-3-0-1-1-1-2-1-0-1-3: Adjust Your Strategy

root-0-0-2-3-0-1-1-1-2-1-0-1-3-2: Anticipate Player Behavior
root-0-0-2-3-0-1-1-1-2-1-0-1-3-4: Risk Management
root-0-0-2-3-0-1-1-1-2-1-0-1-4: Collaborate with Consistent Contributors
root-0-0-2-3-0-1-1-1-2-1-0-1-4-0: Identify Reliable Contributors
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1: Establish Communication
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1-2: Observe Behavioral Patterns
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1-3: Formulate a Joint Strategy
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1-3-1: Optimal Contribution Levels
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1-3-2: Establish Communication
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1-3-3: Adaptation and Flexibility
root-0-0-2-3-0-1-1-1-2-1-0-1-4-1-3-4: Trust and Collaboration
root-0-0-2-3-0-1-1-1-2-1-0-1-4-3: Monitor Consistency
root-0-0-2-3-0-1-1-1-2-1-0-4: Communication and Collaboration
root-0-0-2-3-0-1-1-1-2-1-0-4-2: Encourage Consistency
root-0-0-2-3-0-1-1-1-2-1-0-4-3: Form Alliances
root-0-0-2-3-0-1-1-1-2-1-0-4-3-1: Establish Communication
root-0-0-2-3-0-1-1-1-2-1-0-4-3-2: Coordinate Contribution Efforts
root-0-0-2-3-0-1-1-1-2-1-0-4-3-3: Build Trust and Reliability
root-0-0-2-3-0-1-1-1-2-1-0-4-4: Monitor and Adapt
root-0-0-2-3-0-1-1-1-2-1-2: Communicate and Negotiate
root-0-0-2-3-0-1-1-1-2-1-2-0: Analyze Contribution Patterns
root-0-0-2-3-0-1-1-1-2-1-2-3: Monitor Trustworthiness
root-0-0-2-3-0-1-1-1-2-1-2-4: Adapt to Changing Dynamics
root-0-0-2-3-0-1-1-1-2-1-2-4-1: Form Alliances
root-0-0-2-3-0-1-1-1-2-1-2-4-4: Long-term Planning
root-0-0-2-3-0-1-1-1-2-1-2-4-4-0: Assess the Current Trend
root-0-0-2-3-0-1-1-1-2-1-2-4-4-4: Flexibility in Strategy
root-0-0-2-3-0-1-1-1-2-1-2-4-4-5: Consistency in Contributions
root-0-0-2-3-0-1-1-1-2-1-4: Build a Reputation
root-0-0-2-3-0-1-1-1-2-1-4-2: Observation and Adaptation
root-0-0-2-3-0-1-1-1-2-1-4-4: Communication and Collaboration
root-0-0-2-3-0-1-1-1-2-2: Establish Collaborative Partnerships
root-0-0-2-3-0-1-1-1-2-2-0: Identify Trustworthy Players
root-0-0-2-3-0-1-1-1-2-2-0-2: Consider Long-Term Behavior
root-0-0-2-3-0-1-1-1-2-2-0-2-1: Identify Trustworthy Players

1225
1226
1227
1228
1229
1230
1231
1232
1233
1234
1235
1236
1237
1238
1239
1240
1241
1242
1243
1244
1245
1246
1247
1248
1249
1250
1251
1252
1253
1254
1255
1256
1257
1258
1259
1260
1261
1262
1263
1264
1265
1266
1267
1268
1269
1270
1271
1272
1273
1274
1275
1276
1277
1278
1279
1280
1281
1282
1283
1284
1285
1286
1287
1288
1289
1290
1291
1292
1293
1294

1295	root-0-0-2-3-0-1-1-1-2-2-0-2-3: Adjust	root-0-0-2-3-0-1-4-1: Assess Impact on	1365
1296	Your Strategy	Public Good Payoff	1366
1297	root-0-0-2-3-0-1-1-1-2-2-0-3: Form	root-0-0-2-3-0-1-4-1-1: Evaluate Public	1367
1298	Alliances	Pot Growth	1368
1299	root-0-0-2-3-0-1-1-1-2-2-0-3-1: Assess	root-0-0-2-3-0-1-4-1-3: Identify	1369
1300	Trustworthiness	Collaborative Strategies	1370
1301	root-0-0-2-3-0-1-1-1-2-2-0-3-3: Mutual	root-0-0-2-3-0-1-4-1-4: Predict Future	1371
1302	Benefit	Payoff Trends	1372
1303	root-0-0-2-3-0-1-1-1-2-2-0-3-4: Long-	root-0-0-2-3-0-1-4-2: Compare Individual	1373
1304	Term Collaboration	Gains	1374
1305	root-0-0-2-3-0-1-1-1-2-2-0-4: Monitor	root-0-0-2-3-0-1-4-4: Formulate	1375
1306	Changes	Collaboration Tactics	1376
1307	root-0-0-2-3-0-1-1-1-2-2-1: Initiate	root-0-0-2-3-0-2: Detect Potential	1377
1308	Communication	Competition	1378
1309	root-0-0-2-3-0-1-1-1-2-2-2: Reciprocate	root-0-0-2-3-2: Strategic Adaptation	1379
1310	Trust	root-0-0-2-3-2-0: Analyze Other Players'	1380
1311	root-0-0-2-3-0-1-1-1-2-2-4: Adaptability	Contributions	1381
1312	root-0-0-2-3-0-1-1-1-2-2-4-0: Assess	root-0-0-2-3-2-4: Flexibility in	1382
1313	Other Players' Contributions	Decision Making	1383
1314	root-0-0-2-3-0-1-1-1-2-2-4-2: Identify	root-0-0-2-3-2-4-1: Adjust Contribution	1384
1315	Potential Alliances	Based on Public Pot Size	1385
1316	root-0-0-2-3-0-1-1-1-4: Long-term	root-0-0-2-3-2-4-2: Balance Risk and	1386
1317	Planning	Reward	1387
1318	root-0-0-2-3-0-1-1-1-4-2: Encourage	root-0-0-2-3-2-4-2-0: Assess the Current	1388
1319	Cooperative Behavior	Token Balance	1389
1320	root-0-0-2-3-0-1-1-1-4-2-0: Establish	root-0-0-2-3-2-4-2-2: Adapt Contribution	1390
1321	Trust	Strategy	1391
1322	root-0-0-2-3-0-1-1-1-4-2-1: Strategic	root-0-0-2-3-2-4-2-4: Observe Patterns	1392
1323	Communication	root-0-0-2-3-3: Long-term Planning	1393
1324	root-0-0-2-3-0-1-1-1-4-2-1-2: Highlight	root-0-0-2-3-4: Risk Assessment	1394
1325	Long-Term Benefits	root-0-0-2-3-4-0: Analyze Previous	1395
1326	root-0-0-2-3-0-1-1-1-4-2-1-3: Negotiate	Rounds	1396
1327	Contribution Strategies	root-0-0-2-3-4-0-1: Gain Assessment	1397
1328	root-0-0-2-3-0-1-1-1-4-2-1-4: Foster	root-0-0-2-3-4-0-2: Competitive	1398
1329	Trust and Collaboration	Strategies	1399
1330	root-0-0-2-3-0-1-1-1-4-2-2: Highlight	root-0-0-2-3-4-0-3: Collaboration	1400
1331	Mutual Gains	Opportunities	1401
1332	root-0-0-2-3-0-1-1-1-4-2-3: Foster	root-0-0-2-3-4-2: Assess Potential	1402
1333	Collaboration	Losses	1403
1334	root-0-0-2-3-0-1-1-1-4-2-4: Long-Term	root-0-0-2-3-4-4: Long-term Planning	1404
1335	Perspective	root-0-0-2-4: Long-term Planning	1405
1336	root-0-0-2-3-0-1-1-1-4-3: Monitor and	root-0-0-2-4-0: Monitor Token Balance	1406
1337	Adapt	root-0-0-2-4-0-0: Analyze Contribution	1407
1338	root-0-0-2-3-0-1-1-1-4-3-1: Build	Impact	1408
1339	Sustainable Partnerships	root-0-0-2-4-0-0-2: Strategy	1409
1340	root-0-0-2-3-0-1-1-1-4-3-3: Strategic	Effectiveness	1410
1341	Observation	root-0-0-2-4-0-0-2-0: Contribution	1411
1342	root-0-0-2-3-0-1-1-1-4-3-4: Long-term	Analysis	1412
1343	Adaptation	root-0-0-2-4-0-0-2-0-2: Identify rounds	1413
1344	root-0-0-2-3-0-1-1-1-4-4: Evaluate Long-	with lower gain than expected and	1414
1345	Term Gains	analyze potential reasons	1415
1346	root-0-0-2-3-0-1-1-1-4-4-2: Monitor	root-0-0-2-4-0-0-2-0-3: Experiment with	1416
1347	Contribution Trends	different contribution amounts in future	1417
1348	root-0-0-2-3-0-1-1-2: Monitor Changes in	rounds	1418
1349	Contributions	root-0-0-2-4-4: Risk Management	1419
1350	root-0-0-2-3-0-1-1-2-2: Form	root-0-0-2-4-4-0: Assess Potential Gains	1420
1351	Partnerships	root-0-0-2-4-4-0-0: Analyze Contribution	1421
1352	root-0-0-2-3-0-1-1-2-2-1: Establish	Impact	1422
1353	Communication	root-0-0-2-4-4-1: Balance Contribution	1423
1354	root-0-0-2-3-0-1-1-2-2-2: Form Strategic	root-0-0-2-4-4-3: Long-term Planning	1424
1355	Alliances	root-0-0-2-4-4-4: Flexibility in	1425
1356	root-0-0-2-3-0-1-1-2-2-4: Maximize	Contributions	1426
1357	Collective Gain	root-0-3: Adaptability	1427
1358	root-0-0-2-3-0-1-1-2-3: Anticipate	root-0-3-2: Observation and Prediction	1428
1359	Changes	root-0-3-2-1: Predict Potential	1429
1360	root-0-0-2-3-0-1-1-2-4: Evaluate Risk-	Strategies	1430
1361	Reward Ratio	root-0-3-2-1-0: Player 1	1431
1362	root-0-0-2-3-0-1-3: Build Trust and	root-0-3-2-1-1: Player 2	1432
1363	Cooperation	root-0-3-2-1-2: Player 3	1433
1364	root-0-0-2-3-0-1-4: Monitor Results	root-0-3-2-2: Adjust Your Strategy	1434

1435	root-0-3-2-4: Stay Flexible	Guess 2/3 of the Average:	1505
1436	root-0-3-3: Risk Assessment		1506
1437	root-0-3-3-1: Consider Contribution	root: Choose a number that you believe	1507
1438	Variability	will be closest to 2/3 of the average of	1508
1439	root-0-3-3-1-1: Predict Potential	all numbers chosen by players,	1509
1440	Contributions	including your selection	1510
1441	root-0-3-4: Long-term Adaptation	root-0: Observation	1511
1442	root-0-3-4-2: Flexibility in	root-0-0: Analyze Trends	1512
1443	Contribution	root-0-0-1: Evaluate Deviations	1513
1444	root-0-3-4-2-2: Balance Short-term Gains	root-0-0-1-3: Stay Informed	1514
1445	and Long-term Goal	root-0-0-1-3-3: Flexibility in Decision-	1515
1446	root-0-4: Risk Assessment	Making	1516
1447	root-0-4-0: Analyze Previous Rounds	root-0-0-1-3-3-1: Adapt to Changing	1517
1448	root-0-4-0-1: Risk Assessment	Dynamics	1518
1449	root-0-4-0-1-0: Analyze Previous Rounds	root-0-0-1-3-3-1-3: Consider Risk-Reward	1519
1450	root-0-4-0-1-1: Consider Variability	root-0-0-1-3-3-2: Consider Risk-Reward	1520
1451	root-0-4-0-1-3: Risk Tolerance	Tradeoff	1521
1452	root-0-4-0-1-4: Strategic Adjustment	root-0-0-1-3-3-2-3: Adapt to Changing	1522
1453	root-0-4-0-3: Strategic Planning	Circumstances	1523
1454	root-0-4-4: Adaptation	root-0-0-1-3-3-2-3-3: Strategic	1524
1455	root-1: Strategic Decision Making	Observation	1525
1456	root-1-0: Analyze Other Players'	root-0-0-1-3-3-2-3-3-1: Consider Recent	1526
1457	Contributions	Rounds	1527
1458	root-1-0-3: Consider Overall Game	root-0-0-1-3-3-2-3-3-2: Identify	1528
1459	Dynamics	Outliers	1529
1460	root-1-0-3-1: Assess Token Distribution	root-0-0-1-3-3-2-3-3-3: Predict	1530
1461	root-1-1: Consider Potential Payoff	Potential Average	1531
1462	root-1-1-2: Risk Assessment	root-0-0-1-3-3-2-3-4: Risk Assessment	1532
1463	root-1-1-2-0: Analyze Previous Rounds	root-0-0-1-3-3-4: Balance Consistency	1533
1464	root-1-1-2-0-0: Contribution Level	and Adaptability	1534
1465	Analysis	root-0-0-1-3-4: Strategic Observation	1535
1466	root-1-1-2-0-2: Trend Identification	root-0-0-1-3-4-0: Analyze Winning	1536
1467	root-1-1-2-0-2-0: Consider the overall	Numbers	1537
1468	game dynamics	root-0-0-1-3-4-0-1: Identify Common	1538
1469	root-1-1-2-0-2-1: Flexibility in	Numbers	1539
1470	contribution strategies	root-0-0-1-3-4-0-2: Consider the Average	1540
1471	root-1-1-2-0-2-2: Risk management	root-0-0-1-3-4-1: Monitor Average	1541
1472	root-1-1-2-0-2-2-0: Analyze Trends	Numbers	1542
1473	root-1-1-2-0-2-2-2: Diversify	root-0-0-1-3-4-1-2: Consider Previous	1543
1474	Contributions	Results	1544
1475	root-1-1-2-0-2-3: Observation of player	root-0-0-1-3-4-1-4: Adjust Risk	1545
1476	behavior	Tolerance	1546
1477	root-1-1-2-0-3: Risk Assessment	root-0-0-1-3-4-2: Observe Your	1547
1478	root-1-1-2-0-4: Adaptation Strategy	Performance	1548
1479	root-1-1-2-0-4-2: Consider Overall Game	root-0-0-1-3-4-3: Consider Player	1549
1480	Dynamics	Strategies	1550
1481	root-1-1-2-4: Long-term Risk Management	root-0-0-1-3-4-3-0: Analyze Winning	1551
1482	root-1-1-3: Adapt to Player Behaviors	Strategies	1552
1483	root-1-1-3-2: Strategic Decision Making	root-0-0-1-3-4-3-1: Adaptation	1553
1484	root-1-3: Adapt to Player Behaviors	root-0-0-1-3-4-3-2: Observation	1554
1485	root-1-3-3: Balance Risk and Reward	root-0-0-1-3-4-3-4: Risk Assessment	1555
1486	root-1-5: Flexibility	root-0-1: Identify Outliers	1556
1487	root-1-5-1: Adjust Contribution Based on	root-0-1-0: Analyze Previous Rounds	1557
1488	Public Pot	root-0-1-0-1: Consider Trends	1558
1489	root-1-5-1-0: Analyze Public Pot Size	root-0-1-0-1-0: Consider the decreasing	1559
1490	root-1-5-1-0-2: Monitor Overall Trends	trend in the average number chosen by	1560
1491	root-1-5-1-0-2-2: Compare with Other	players in the previous rounds and	1561
1492	Players	select a number slightly lower than the	1562
1493	root-1-5-1-2: Monitor Overall Token	expected average for the upcoming round	1563
1494	Accumulation	root-0-1-0-1-0-3: Balance Risk and	1564
1495	root-2: Long-term Planning	Reward	1565
1496	root-2-0: Assess Previous Contributions	root-0-1-0-1-0-3-2: Cautious Approach	1566
1497	root-2-0-1: Identify Optimal	root-0-1-0-1-0-3-3: Strategic Thinking	1567
1498	Contribution Levels	root-0-1-0-1-0-3-5: Observation	1568
1499	root-2-0-2: Consider Player Behaviors	root-0-1-0-1-0-4: Monitor Results	1569
1500	root-2-0-3: Adjust Contribution Strategy	root-0-1-0-2: Adjust for Variability	1570
1501	root-2-1: Strategic Contribution	root-0-1-0-2-0: Analyze Previous	1571
1502	root-2-2: Monitor Other Players	Averages	1572
1503		root-0-1-0-2-0-1: Identify Trends	1573
1504		root-0-1-0-2-0-1-2: Consider the Range	1574

1575	root-0-1-0-2-0-2: Consider Outliers	root-0-4-1-3-1: Anticipate Competitors' Choices	1645
1576	root-0-1-0-2-0-2-0: Analyze Previous Outliers		1646
1577		root-0-4-1-3-3: Consider Risk-Reward	1647
1578	root-0-1-0-2-0-2-3: Factor in Player Behavior	root-0-4-3: Stay Informed	1648
1579		root-0-4-4: Utilize Strategic Thinking	1649
1580	root-0-1-0-2-0-2-3-1: Identify Player Tendencies	root-1: Strategic Thinking	1650
1581		root-1-2: Calculating 2/3 of the Average	1651
1582	root-0-1-0-2-0-2-3-2: Adjust Number Selection	root-1-3: Strategic Number Selection	1652
1583		root-1-4: Adaptation and Flexibility	1653
1584	root-0-1-0-2-1: Consider Conservative Approach	root-1-4-2: Evaluate Your Own Strategy	1654
1585		root-1-4-4: Stay Informed	1655
1586	root-0-1-0-2-1-1: Identify Central Tendency	root-1-4-5: Strategic Variation	1656
1587		root-2: Risk Assessment	1657
1588	root-0-1-0-2-1-2: Avoid Extreme Outliers	root-2-1: Consider Variability	1658
1589	root-0-1-0-2-1-3: Consider Stability	root-2-3: Assess Risk Tolerance	1659
1590	root-0-1-0-2-1-4: Balance Risk and Reward	root-2-4: Anticipate Strategic Play	1660
1591		root-3: Adaptation	1661
1592	root-0-1-0-2-1-4-1: Consider the Current Average	root-3-3: Risk Assessment	1662
1593		root-3-3-1: Consider the Range	1663
1594	root-0-1-0-2-1-4-2: Assess Your Position	root-3-3-4: Utilize Previous Experience	1664
1595	root-0-1-0-2-1-4-4: Adapt to the Game Dynamics	root-4: Long-term Planning	1665
1596		root-4-2: Strategic Adjustment	1666
1597	root-0-1-0-2-1-4-5: Stay Informed	root-4-4: Risk Assessment	1667
1598	root-0-1-0-2-2: Evaluate Trends	root-4-4-1: Consider Variability	1668
1599	root-0-1-0-2-4: Adapt to Changing Dynamics	root-4-4-2: Evaluate Your Performance	1669
1600			1670
1601	root-0-1-0-2-4-1: Flexibility in Number Selection		1671
1602		Auction Arena:	1672
1603	root-0-1-0-2-4-2: Consider Outliers		1673
1604	root-0-1-0-2-4-4: Risk Assessment	root: secure the highest profit at the end of this auction, compared to all other bidders	1674
1605	root-0-1-1: Consider Potential Influences	root-0: Efficiently allocate budget	1675
1606		root-0-0: Prioritize items with a higher difference between your estimated value and the starting price	1676
1607	root-0-1-2: Predict Potential Outliers	root-0-0-1: Consider the competition	1677
1608	root-0-1-2-0: Analyze the Trend	root-0-0-1-1: Identify Weaknesses	1678
1609	root-0-1-3: Adjust Your Strategy	root-0-0-1-1-1: Monitor Budget Utilization	1679
1610	root-0-1-3-1: Consider the Trend	root-0-0-1-1-1-1: Strategically Allocate Bids	1680
1611	root-0-1-3-1-1: Adjust Strategy	root-0-0-1-1-1-1-2: Monitor Competitor Bids	1681
1612	root-0-1-3-1-2: Stay Vigilant	root-0-0-1-1-1-1-2-1: Strategic Allocation of Bids	1682
1613	root-0-1-3-2: Balance Risk and Reward	root-0-0-1-1-1-1-2-1-1: Focus on Items with Less Interest	1683
1614	root-0-1-3-2-1: Consider the Impact of Outliers	root-0-0-1-1-1-1-2-1-2: Monitor Potential Withdrawals	1684
1615		root-0-0-1-1-1-1-2-2: Budget Conservation	1685
1616	root-0-1-3-2-1-0: Analyze Previous Rounds	root-0-0-1-1-1-1-4: Maintain Flexibility	1686
1617		root-0-0-1-1-2: Assess Risk-Taking Behavior	1687
1618	root-0-1-3-2-1-1: Adjust Strategy	root-0-0-1-1-2-1: Identify Weaknesses	1688
1619	root-0-1-3-2-1-2: Monitor Extreme Numbers	root-0-0-1-1-2-1-0: Analyze Bidding Patterns	1689
1620		root-0-0-1-1-2-1-3: Monitor Remaining Items	1690
1621	root-0-1-3-2-1-4: Stay Flexible	root-0-0-1-1-2-3: Budget Management	1691
1622	root-0-1-3-2-4: Stay Informed	root-0-0-1-1-3: Identify Overestimation	1692
1623	root-0-1-3-3: Adapt to Competitors	root-0-0-1-1-4: Exploit Predictable Behavior	1693
1624	root-0-1-3-3-1: Balance Risk and Reward	root-0-0-1-2: Formulate Counter-Strategies	1694
1625	root-0-1-3-3-2: Anticipate Competitors' Choices	root-0-0-1-2-4: Psychological Tactics	1695
1626		root-0-0-1-3: Adaptability	1696
1627	root-0-1-3-3-2-4: Flexibility	root-0-0-1-3-1: Adjust Bidding Strategy	1697
1628	root-0-1-3-3-4: Strategic Risk-Taking	root-0-0-1-3-4: Evaluate Risk-Reward	1698
1629	root-0-1-3-3-4-2: Consider the Range		1699
1630	root-0-1-3-3-4-3: Balance Consistency and Differentiation		1700
1631	root-0-1-3-3-4-4: Adapt Based on Previous Outcomes		1701
1632	root-0-2: Consider Player Behavior		1702
1633	root-0-2-1: Adjust Based on Averages		1703
1634	root-0-2-3: Stay Flexible		1704
1635	root-0-2-3-2: Evaluate Your Position		1705
1636	root-0-2-3-3: Monitor Player Behaviors		1706
1637	root-0-3: Factor in Previous Results		1707
1638	root-0-3-1: Consider Trend		1708
1639	root-0-4: Adjust Strategy		1709
1640	root-0-4-1: Consider Your Competitors		1710
1641	root-0-4-1-1: Adjust for Biases		1711
1642	root-0-4-1-3: Use Game Theory		1712
1643			1713
1644			1714

1715	Ratio	negotiation, regardless of your own	1785
1716	root-0-0-1-5: Information Utilization	profit.	1786
1717	root-0-0-1-5-0: Analyze Bidders' Behavior	root-0: Maximize the number of items you receive	1787
1718			1788
1719	root-0-0-1-5-1: Adjust Bidding Strategy	root-0-0: Evaluate the value of each item	1789
1720	root-0-0-1-5-1-0: Analyze Previous Bidding Patterns		1790
1721		root-0-1: Consider trade-offs	1791
1722	root-0-0-1-5-1-0-1: Target Items with Lower Competition	root-0-2: Seek compromise	1792
1723		root-0-3: Communicate effectively	1793
1724	root-0-0-1-5-1-0-3: Evaluate True Values	root-0-4: Be flexible	1794
1725	root-0-0-1-5-1-2: Evaluate Profit Margins	root-1: Prioritize high-value items	1795
1726		root-1-0: Assess the value of each item	1796
1727	root-0-0-1-5-1-3: Identify High-Value Items	root-1-1: Consider trade-offs	1797
1728		root-1-2: Negotiate for high-value items	1798
1729	root-0-0-1-5-1-6: Adapt to True Values	root-1-3: Be open to compromise	1799
1730	root-0-1: Monitor the bidding behavior of other bidders	root-1-4: Communicate the reasoning behind your prioritization	1800
1731		root-2: Ensure fair distribution	1801
1732	root-0-1-2: Strategic Bidding	root-2-0: Consider the value of each item	1802
1733	root-0-1-2-5: Stay Informed		1803
1734	root-0-3: Be prepared to adjust your estimated value	root-2-1: Propose a balanced allocation	1804
1735		root-2-2: Be open to compromise	1805
1736	root-0-4: Aim for a balance between winning bids and maximizing profit	root-2-3: Communicate the reasoning behind your proposal	1806
1737		root-2-4: Seek mutual agreement	1807
1738	root-1: Accurately estimate item values	root-3: Maintain a cooperative and communicative approach	1808
1739	root-1-0: Research	root-3-0: Clarify interests and priorities	1809
1740	root-1-1: Analyze Previous Auctions		1810
1741	root-1-1-1: Analyze Market Trends	root-3-1: Seek common ground	1811
1742	root-1-1-1-0: Research Market Demand	root-3-2: Explore trade-offs	1812
1743	root-1-1-1-1: Consider Seasonality	root-3-3: Remain open to creative solutions	1813
1744	root-1-1-1-2: Economic Conditions	root-3-4: Maintain a positive and respectful tone	1814
1745	root-1-1-2: Adjust Estimated Values	root-4: Adapt and adjust strategies	1815
1746	root-1-2: Consider Item Condition	root-4-0: Understand Bob's priorities	1816
1747	root-1-3: Adjust Estimations	root-4-2: Propose alternative allocations	1817
1748	root-1-3-1: Consider True Value	root-4-3: Maintain open communication	1818
1749	root-1-3-4: Adapt to Competition	root-4-4: Be willing to compromise	1819
1750	root-1-4: Budget Management		1820
1751	root-1-4-1: Risk Assessment		1821
1752	root-1-4-2: Prioritize High-Value Items		1822
1753	root-1-4-2-0: Assess Remaining Budget		1823
1754	root-1-4-2-3: Monitor Competing Bidders		1824
1755	root-1-5: Risk Assessment		1825
1756	root-2: Strategic bidding		
1757	root-2-0: Budget Management		
1758	root-2-1: Estimated Value Comparison		
1759	root-2-2: Observation of Competitors		
1760	root-2-3: Risk Assessment		
1761	root-2-4: Strategic Withdrawal		
1762	root-2-4-0: Assess Potential Profit Margin		
1763			
1764	root-2-4-5: Long-term Profit Maximization		
1765			
1766	root-3: Risk management		
1767	root-3-1: Budget Allocation		
1768	root-3-2: Competitive Analysis		
1769	root-3-2-1: Assess Remaining Competitors		
1770	root-3-2-2: Estimate Competitors' Valuation		
1771			
1772	root-3-3: Flexibility in Bidding		
1773	root-3-5: Information Gathering		
1774	root-3-5-1: Refine risk assessment		
1775	root-3-5-4: Anticipate competition		
1776	root-3-5-5: Adapt bidding strategy		
1777	root-4: Adaptability		
1778	root-4-4: Risk Management		
1779	root-4-6: Adapt to Market Dynamics		
1780			
1781	DealOrNotDeal		
1782			
1783	root: minimize the profit gap between yourself and your partner in this		
1784			