

LEVERAGING IMITATION LEARNING AND LLMs FOR EFFICIENT HIERARCHICAL REINFORCEMENT LEARNING

Anonymous authors

Paper under double-blind review

ABSTRACT

In this paper, we introduce an innovative framework that combines Hierarchical Reinforcement Learning (HRL) with Large Language Models (LLMs) to tackle the challenges of complex, sparse-reward environments. A key contribution of our approach is the emphasis on imitation learning during the early training stages, where the LLM plays a crucial role in guiding the agent by providing high-level decision-making strategies. This early-stage imitation learning significantly accelerates the agent’s understanding of task structure, reducing the time needed to adapt to new environments. By leveraging the LLM’s ability to generate abstract representations of the environment, the agent can efficiently explore potential strategies, even in tasks with high-dimensional state spaces and delayed rewards. Our method utilizes the LLM to assist action sampling via a dynamic annealing strategy and aids the policy learning process through an LLM-based policy and value regularizer. This approach reduces computational costs and enhances the agent’s learning process. Experimental results across three environments—MiniGrid, NetHack, and Crafter—demonstrate that our method significantly outperforms baseline LLM-based HRL algorithms in terms of training speed and success rates. The imitation learning phase proves critical in enabling the agent to adapt quickly and perform efficiently, highlighting the potential of integrating LLMs into HRL for complex tasks.

1 INTRODUCTION

Reward sparsity is a persistent challenge in the early stages of exploration for reinforcement learning (RL) environments. As these environments grow more complex, the difficulty for agents to encounter rewards during initial exploration increases significantly. Researchers have continuously worked on addressing reward sparsity to efficiently train agents (Vecerik et al., 2017; Hare, 2019). Hierarchical reinforcement learning (HRL) offers a promising approach to addressing these issues. In HRL, the decision-making problem is decomposed into two levels: high-level decision making, referred to as an *option*, and low-level decision making, referred to as an *action*. Experimental evidence (Kulkarni et al., 2016; Nachum et al., 2018b) suggests that HRL can address challenges that traditional RL algorithms struggle with, demonstrating superior performance in general environments. However, despite HRL’s potential, most approaches rely on predefined options and often require pre-training of the option networks. This makes HRL notoriously difficult to implement and tune effectively.

LLMs have shown potential in mitigating some of these challenges. Recently, LLMs have demonstrated their versatility across many domains, including natural language processing, code generation, and decision-making tasks such as game playing and intelligent question-answering (Du et al., 2023). A series of recent works (Zhou et al., 2024) suggest that LLMs can be leveraged as high-level decision-makers to enhance the performance of RL agents. While promising, it is well-known that LLM-based approaches are resource-intensive, which makes them less ideal compared to traditional RL methods that are computationally lighter. This leads to the following question:

What is the most effective strategy to accelerate RL with LLMs?

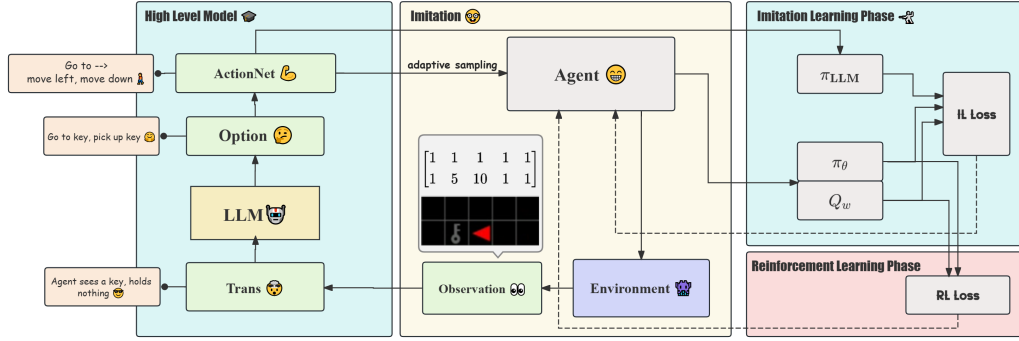


Figure 1: Algorithm framework of IHAC. Our framework is a two-phase algorithm designed to improve the learning efficiency of reinforcement learning agents by incorporating high-level guidance from LLM. In the first phase, imitation learning is used, where the agent benefits from LLM-generated options to accelerate its exploration and learning process. The second phase transitions to standard reinforcement learning, where the agent refines its policy and value networks, initially trained with the help of the LLM, to further optimize its decision-making abilities. By balancing imitation learning with reinforcement learning, IHAC achieves a more efficient learning process, especially in complex and long-term environments with sparse rewards.

To address this question, we propose the Imitation Hierarchical Actor-Critic (IHAC) framework, which accelerates the RL process in a hierarchical RL approach by incorporating high-level instructions from LLMs. The overview of IHAC is displayed in Figure 1. Our contributions are listed as follows.

- We propose a novel approach to imitation learning by framing the decision-making process as a hierarchical RL problem, where high-level options represent overarching instructions that an agent can follow to enhance decision-making. Our key contribution lies in the two-phase design of the proposed IHAC framework, which effectively balances imitation learning and reinforcement learning. In the first phase, IHAC leverages an external LLM to generate a higher-policy action distribution, which serves as input to both the actor and critic for simultaneous imitation learning. This ensures that the LLM’s guidance is maximally utilized, enabling the agent to learn both action selection and value estimation early on, when its interaction with the environment is limited. In the second phase, IHAC transitions to a standard RL algorithm (e.g., PPO) to fine-tune the policy further. This design significantly accelerates learning while reducing reliance on LLM in later stages, achieving computational efficiency without sacrificing performance.
- To fully harness the power of the LLM, we introduce an adaptive sampling strategy that combines input from both the RL agent and the LLM to derive better actions during the imitation learning phase. Additionally, we propose an adaptive policy training strategy that facilitates a more precise approximation of the agent’s policy with the help of the LLM. Both adaptive strategies help us balance the LLM’s guidance with the RL agent’s learning potential.
- For our experiments, we test our algorithm on several standard hierarchical RL benchmarks, such as MiniGrid Chevalier-Boisvert et al. (2023). Compared to existing baselines, our algorithm demonstrates superior performance and greater efficiency, particularly in terms of token usage, due to the incorporation of the LLM and our adaptive algorithm design.

2 RELATED WORKS

2.1 HIERARCHICAL REINFORCEMENT LEARNING (HRL)

The emergence of hierarchical reinforcement learning (HRL) has proven to be beneficial for addressing complex, large-scale problems and sparse, delayed rewards (Nachum et al., 2019). By introducing sub-goals and training multi-level structures, HRL effectively enhances exploration. However, a key challenge in HRL lies in establishing high-quality hierarchical structures to improve training

efficiency. One approach to establishing hierarchical structures involves manually setting them based on tasks, such as graph-guided reinforcement learning (Lee et al., 2022; Gieselmann & Pokorny, 2021), or by artificially setting sub-goals (Florensa et al., 2017; Tessler et al., 2017). The intervention of human priors renders the model non-generalizable. Additionally, there have been proposals to allow agents to autonomously learn sub-tasks. Some achieve this by associating the sub-task space with the current state (Zhang et al., 2020; Nachum et al., 2018a; Vezhnevets et al., 2017), while others consider predicting the next sub-goal during the training process (Pitis et al., 2020).

To enhance training efficiency, some researchers have tackled this issue by constraining the tasks of higher-level agents, limiting the state space of sub-tasks to the adjacent range of the current state (Zhang et al., 2020). Luo et al. (2023) proposed introducing attention rewards to enable higher-level agents to focus more on the environment. Others build a causality-driven hierarchical reinforcement learning framework, leveraging a causality-driven discovery instead of a randomness-driven exploration (Hu et al., 2023). *Hierarchical in-Context Reinforcement Learning* (HCRL) further integrates in-context learning with HRL frameworks to dynamically generate sub-goals based on ongoing task progress. This framework, through reflection and modularity, allows for correction of sub-task errors and improves task efficiency (Sun et al., 2024). *LLM Augmented Hierarchical Agents* have also leveraged LLMs to improve high-level decision-making, leading to better performance in long-horizon tasks (Prakash et al., 2024).

2.2 LLM AGENT

Previous research has demonstrated that language can facilitate the construction of abstract representations of both the environment and goals, especially in scenarios involving high-dimensional state spaces and long planning horizons (Lin et al., 2023; Andreas et al., 2017; Akakzia et al., 2020; Mirchandani et al., 2021). For example, Wu et al. (2019) showed that language could assist in learning complex tasks, even when naive goal representations fail. By employing language-guided hindsight goal relabeling, their approach achieved significant performance improvements. Jiang et al. (2019) further explored the use of language in goal specification, demonstrating that high-level policies could achieve their objectives by composing sub-policies guided by language. These studies underscore the utility of language as a compositional tool in RL, enabling more effective learning and generalization in complex, temporally-extended tasks.

Language-assisted reinforcement learning provides agents with a high-level understanding of tasks and environments. LLMs, with their powerful language processing capabilities, further expand this field by offering more nuanced task guidance, abstract representations, and decision support (Kwon et al., 2023; Du et al., 2023; Hu et al., 2023). Yao et al. (2023) introduced ReAct, a method where the LLM generates “thoughts” to address problems based on observations. Building on ReAct, Shinn et al. (2024) developed Reflexion, which leverages a few-shot verbal feedback approach to improve decision-making abilities. Ahn et al. (2022) provided another viewpoint, endowing LLMs with foundational “skills” along with corresponding value functions and affordance functions. With LLMs’ assistance, each action selection is comprehensively considered, integrating the reward (value) of the current action and its feasibility (affordance). Nottingham et al. (2023) describe how to leverage Brief Language Inputs for Decision-making Responses (BLINDER) to condense the input text to LLMs, effectively reducing the cost of invoking LLMs by approximately sixfold. Building on these foundational studies, Li et al. (2023) proposed Interactive Task Planning (ITP), a framework that leverages LLMs for dynamic task planning and replanning in robotic systems. By integrating high-level planning with low-level skill execution, ITP enables robust adaptation to user feedback and novel tasks without the need for task-specific training or extensive prompt engineering.

2.3 IMITATION LEARNING

In many previous studies, imitation has become a strategy to develop the performance of reinforcement learning.

Introduced by Price & Boutillier (2003), the student agent can observe the state transitions induced by the mentor’s actions and use the information gleaned from these observations to update the estimate of the value of its own states and actions, thus accelerating the process of RL training.

