

SOLVING THE GRANULARITY MISMATCH: HIERARCHICAL PREFERENCE LEARNING FOR LONG- HORIZON LLM AGENTS

Heyang Gao^{1,*}, Zexu Sun^{2,*}, Erxue Min², Hengyi Cai², Shuaiqiang Wang², Dawei Yin², Xu Chen^{1,†}

¹ Gaoling School of Artificial Intelligence, Renmin University of China

² Baidu Inc.

{gaoheyang, xu.chen}@ruc.edu.cn, sunzexu0826@gmail.com

ABSTRACT

Large Language Models (LLMs) as autonomous agents are increasingly tasked with solving complex, long-horizon problems. Aligning these agents via preference-based offline methods like Direct Preference Optimization (DPO) is a promising direction, yet it faces a critical granularity mismatch. Trajectory-level DPO provides a signal that is too coarse for precise credit assignment, while step-level DPO is often too myopic to capture the value of multi-step behaviors. To resolve this challenge, we introduce **Hierarchical Preference Learning (HPL)**, a hierarchical framework that optimizes LLM agents by leveraging preference signals at multiple, synergistic granularities. While HPL incorporates trajectory- and step-level DPO for global and local policy stability, its core innovation lies in group-level preference optimization guided by a dual-layer curriculum. HPL first decomposes expert trajectories into semantically coherent action groups and then generates contrasting suboptimal groups to enable preference learning at a fine-grained, sub-task level. Then, instead of treating all preference pairs equally, HPL introduces a curriculum scheduler that organizes the learning process from simple to complex. This curriculum is structured along two axes: the group length, representing sub-task complexity, and the sample difficulty, defined by the reward gap between preferred and dispreferred action groups. Experiments on three challenging agent benchmarks show that HPL outperforms existing state-of-the-art methods. Our analyses demonstrate that the hierarchical DPO loss effectively integrates preference signals across multiple granularities, while

