

Semantic Vision Priors for Multi-Agent Reinforcement Learning: Improving Unity Tank Battles with Frozen Vision-Language Models

Friday 19 September 2025 10:00 (13 minutes)

Traditional multi-agent reinforcement learning (MARL) struggles in visually rich environments when agents rely solely on raw pixels or low-level features, often leading to poor exploration and cyclic behaviors. In this work, we propose a novel framework that injects semantic vision priors from a frozen vision-language model (VLM) into the RL pipeline to guide both perception and strategy. At each timestep, agents capture camera frames and, together with a concise natural-language prompt, query a pretrained VLM (e.g., CLIP or BLIP-2) to produce a fixed semantic embedding that encodes object identities and spatial relationships. These embeddings are concatenated with standard numeric observations and fed into a lightweight policy network trained with PPO. We further augment training with two enhancements: (1) auxiliary reward shaping, in which VLM-based object detections (e.g., “enemy sighting”) yield small exploration bonuses, and (2) a hierarchical “coach” loop, where the VLM proposes high-level mini-plans every N steps that condition low-level action execution. We outline an experimental evaluation in a Unity tank battle arena comparing (i) baseline MARL, (ii) semantic-obs only, (iii) semantic-obs + reward shaping, and (iv) complete hierarchical coaching. We hypothesize that semantic priors will accelerate learning—escaping aimless circling—and yield superior coordination and win rates within 2 million environment steps. This approach opens new avenues for integrating off-the-shelf VLM knowledge into real-time multi-agent systems.

Authors: RINALDI, Manuel (Universitatea Politehnica Bucuresti); CHIRU, Costin-Gabriel (PUB)

Presenter: RINALDI, Manuel (Universitatea Politehnica Bucuresti)

Session Classification: Doctoral Symposium

Track Classification: Doctoral Symposium