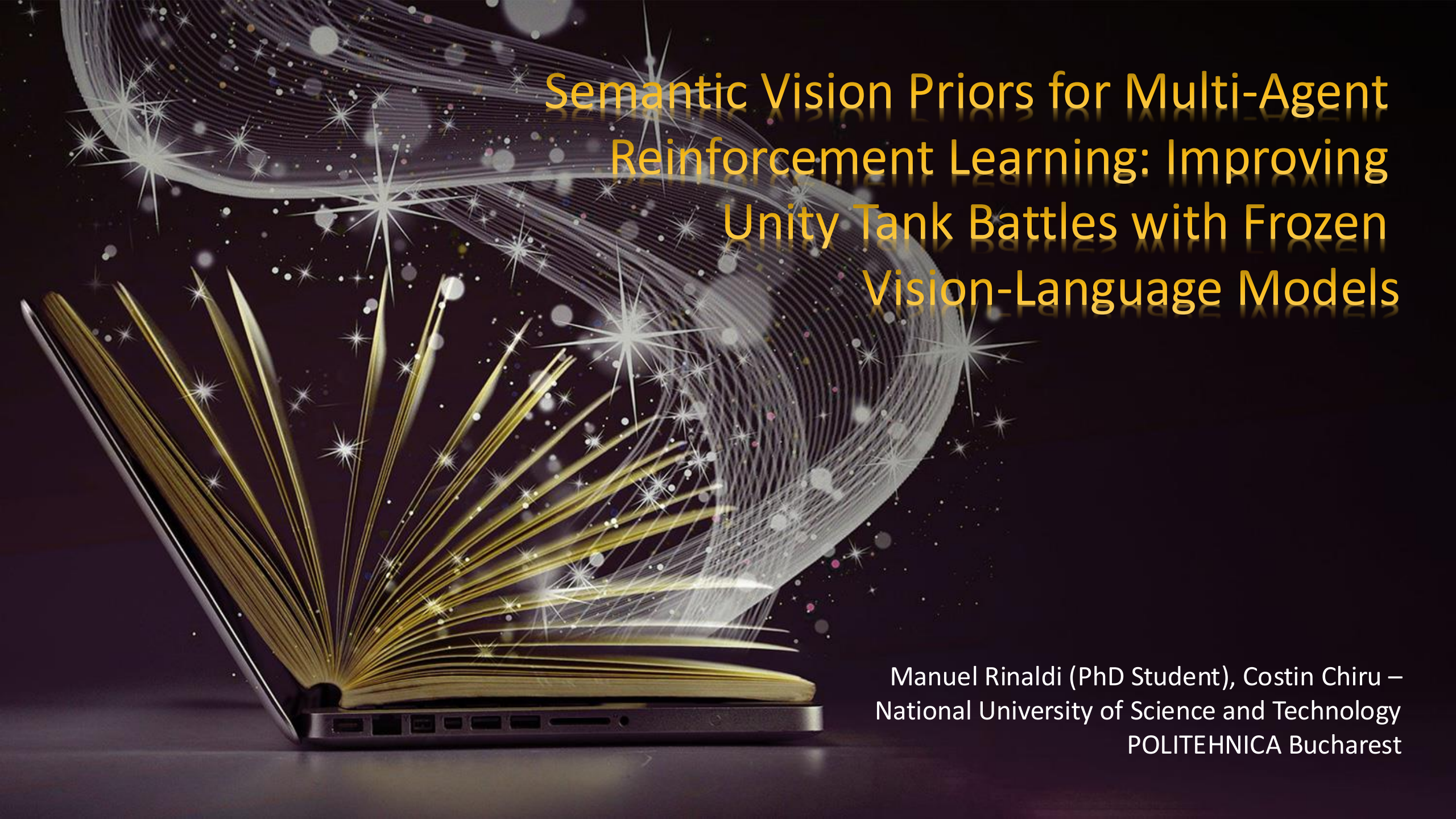


The background of the slide features a dark, textured surface. In the lower-left corner, a laptop is shown from a side-on perspective. The screen of the laptop displays a glowing, golden-yellow book that is open. From the pages of the book, numerous bright, star-like particles and thin, curved lines of light emanate, spreading across the upper half of the image. These particles vary in size and brightness, creating a magical or 'enchanted' atmosphere. The overall color palette is dark, with the primary light source being the golden glow from the book screen.

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

Manuel Rinaldi (PhD Student), Costin Chiru –
National University of Science and Technology
POLITEHNICA Bucharest





Agenda

01 Introduction

Environment, Methodology

02 Related Work

Self-Play Reinforcement Learning, Vision-Language Models in RL,
Semantic Priors in Policy Learning

03 Results

Learning Performance, Self-Play ELO Integration

04 Future Work

Reward Dynamics, Observation modalities, Generalization



Introduction

Environment

Environment



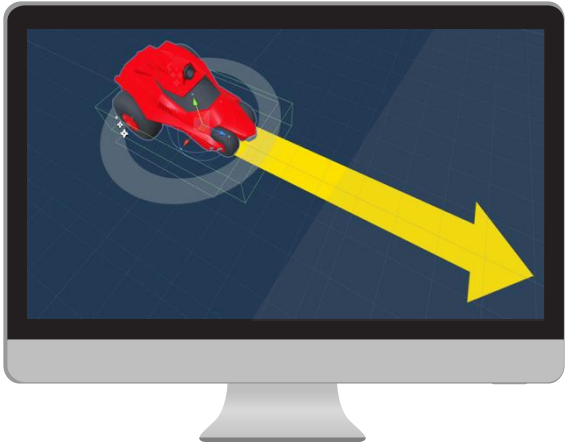
Full Environment



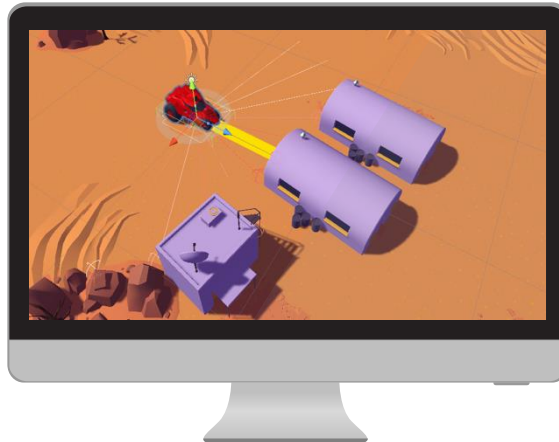
Obstacles



Spawn Point



Tank



RayCast



Tank Camera



Introduction

Methodology

Baseline	Raycast	Semantic Observation
Own Position	Own Position	Own Position
Enemy Position	Raycast Vectors	VLM Embeddings 512
Health	Health	Health
Shooting Recharge	Shooting Recharge	Shooting Recharge
	2M Steps PPO Self Play	





Related Work

Self-Play Reinforcement Learning

Self-Play RL (why it matters)

What it is & why it matters

Self-play has delivered superhuman performance in competitive domains (AlphaGo/AlphaStar/OpenAI Five). Challenge: carrying that success into visually rich 3D remains hard

Common mechanics

PPO self-play with ELO tracking, snapshot pools, opponent rotation; concrete schedule: save every 50k steps, swap 10k, team swap 200k, latest-model ratio 0.5; initial ELO 1200.

Known pitfalls in 3D

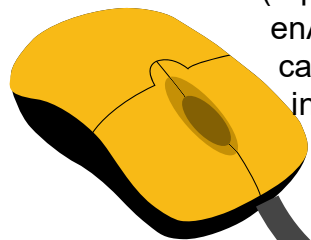
In 3D, self-play can overfit or cycle without careful reward design and observation choices.

Gap my work aims to target

Prior self-play excels in abstract games; your novelty is testing frozen semantic priors inside competitive self-play.

Evidence preview

With the same budget, semantic agents achieve higher ELO despite slower reward; raycasts sit in between; baseline peaks in reward but ELO collapses

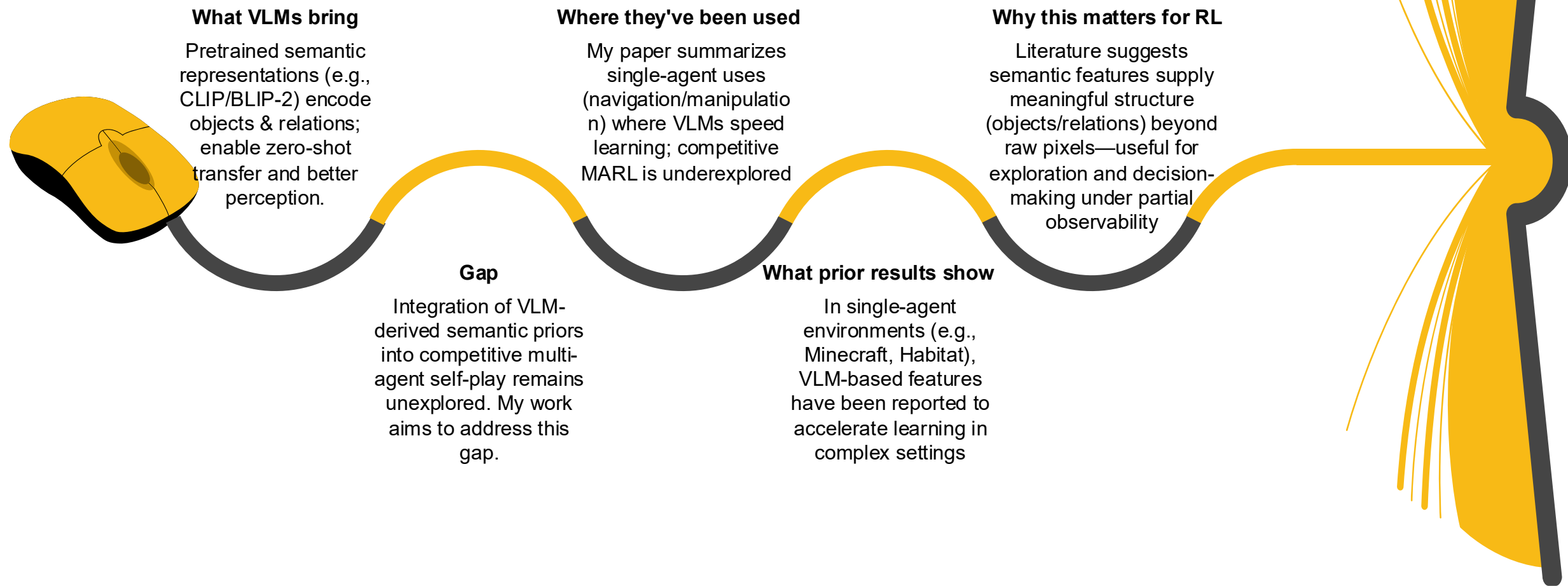




Related Work

Vision-Language Models in RL

Vision-Language Models in RL

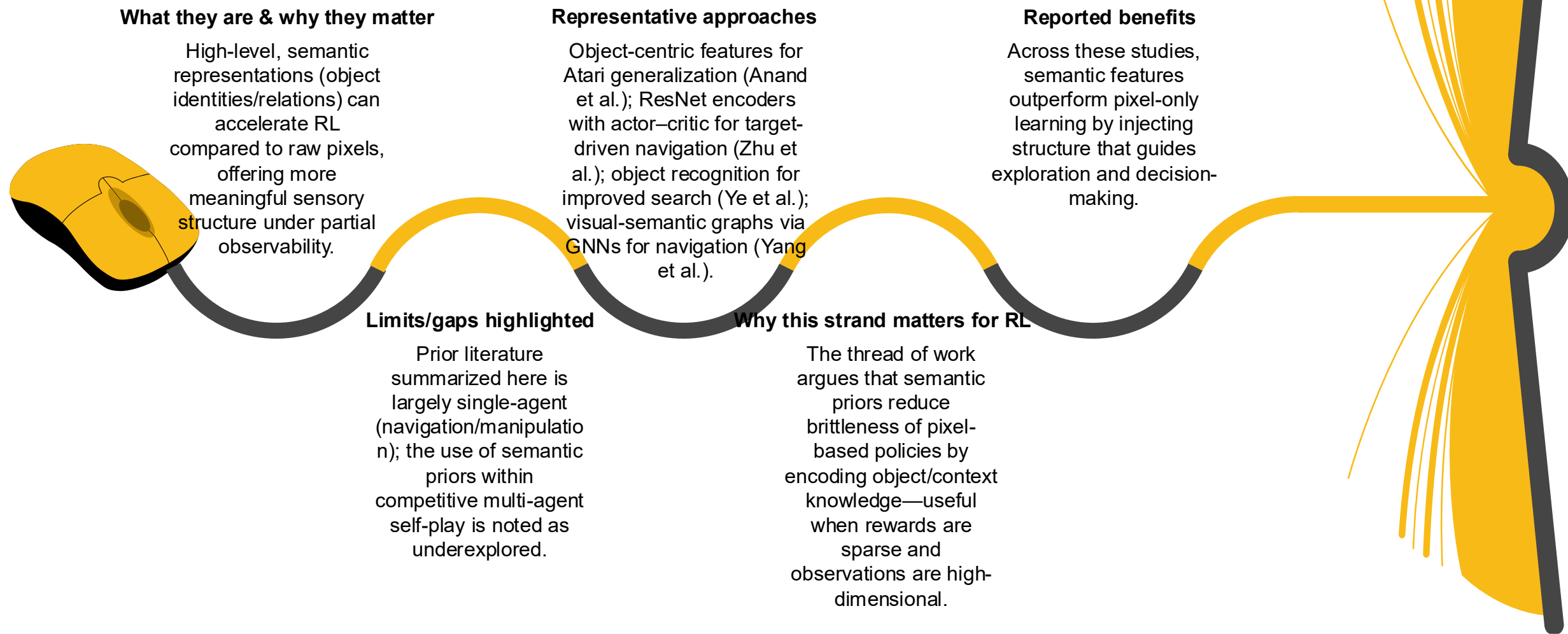




Related Work

Semantic Priors in Policy Learning

Semantic Priors in Policy Learning



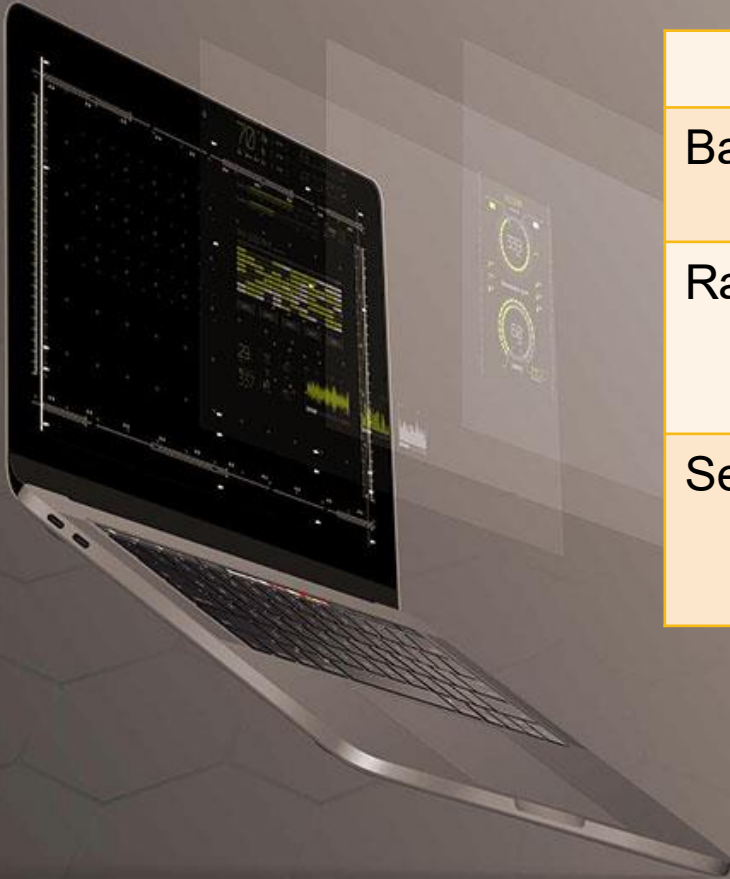


Results

Learning Performance

Learning Performance (2M steps)

Context: 1v1 Unity tanks, PPO self-play, fixed **2M steps**; three observation modalities (Baseline numeric, Raycasts, Semantic BLIP-2).



	Mean Reward	Final Reward	Notes
Baseline	13.53	19.63	Always attack do not explore
Raycast	6.79	8.33	Explore and attack once see the target
Semantic (VLM)	4.49	1.83	Only Explore didn't learn to attack

Overall: All three configurations **improve reward over time**; learning dynamics differ (fast vs moderate vs exploratory).

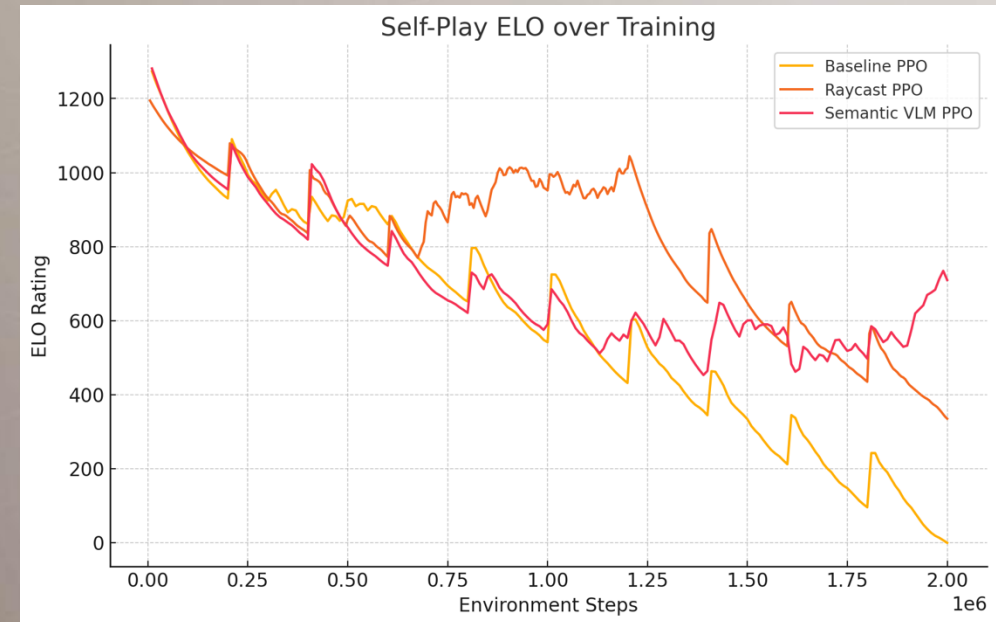


Results

Self-Play ELO

Self-Play ELO (2M steps)

	ELO	Notes
Baseline	-0.45	collapsed from 1273 $\rightarrow$ -0.45
Raycast	334.96	moderate/stabilized
Semantic (VLM)	709.66	highest ELO





Future Work

Reward Dynamics, Observation modalities,
Generalization

Where to go from here?

Rewards

- Experiment using rewards that are either less frequent or more frequent.
- Test rewards that focus more on tactical elements to help the tank accomplish a specific goal.

Semantics & Raycast

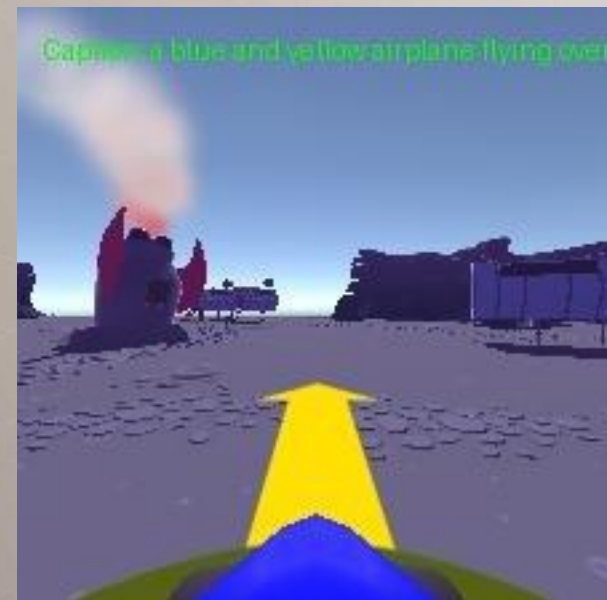
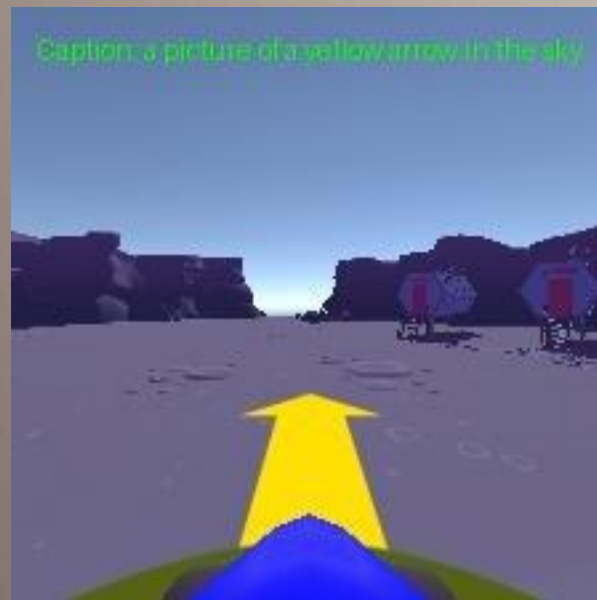
- Try using a larger VLM
- Test a fine-tuned VLM
- Explore additional prompt engineering techniques
- Attempt to generate strategies based on image analysis
- Combine VLM with Raycast for experimentation

Generalization

- Try out the three distinct levels offered
- Utilize Procedural Content Generation to vary the obstacles



Implications of Frozen VLM





THANK YOU