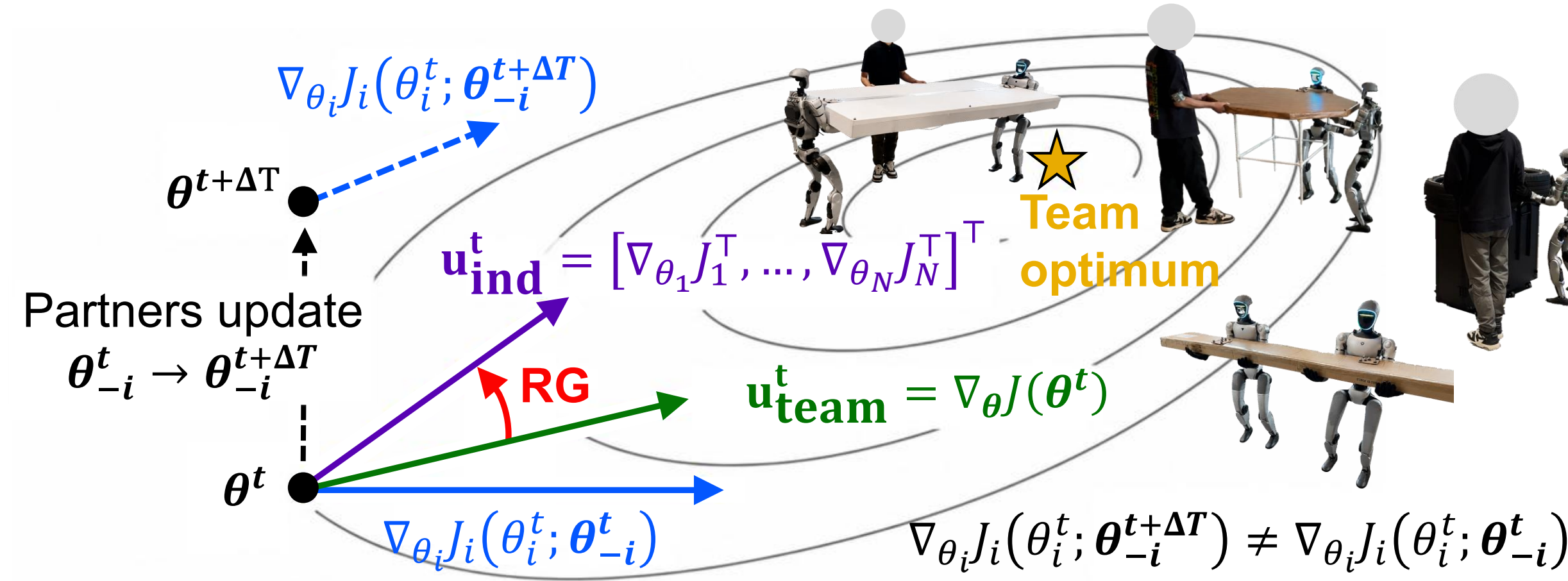


Motivation

Scripted human proxies lead to OOD performance collapse in human-robot collaboration (HRC), motivating multi-agent reinforcement learning (MARL). However, MARL suffers from non-stationarity and a structural rationality gap (RG), as decentralized updates deviate from team-optimal paths.



Problem Formulation

We model HRC as a decentralized POMDP $\mathcal{M} = \langle \mathcal{S}, \mathcal{A}, P, R, \gamma, N, \mathcal{O}, Z \rangle$. To maximize the global return $J(\theta)$, heterogeneous agents utilize decoupled policies $\pi_{\theta_i}(a_{i,t}|o_{i,t})$, operating within a joint parameter $[\theta_1^T, \dots, \theta_N^T]^T \in \mathbb{R}^D$.

Misalignment in Decentralized MARL

Decoupled Learning Dynamics – Interaction between local agent intentions and the global team objective drives learning in CTDE:

- **Independent Rationality Field $\mathbf{u}_{\text{ind}}$** – Assuming other agents are stationary:

$$\mathbf{u}_{\text{ind}}(\theta) \triangleq [\nabla_{\theta_1} J_1^T, \dots, \nabla_{\theta_N} J_N^T]^T \in \mathbb{R}^D$$

- **Team Rationality Field $\mathbf{u}_{\text{team}}$** – True ascent direction of the global team return $J(\theta)$:

$$\mathbf{u}_{\text{team}}(\theta) \triangleq \nabla_{\theta} J(\theta) = \left[\frac{\partial J}{\partial \theta_1}^T, \dots, \frac{\partial J}{\partial \theta_N}^T \right]^T \in \mathbb{R}^D$$

Contact

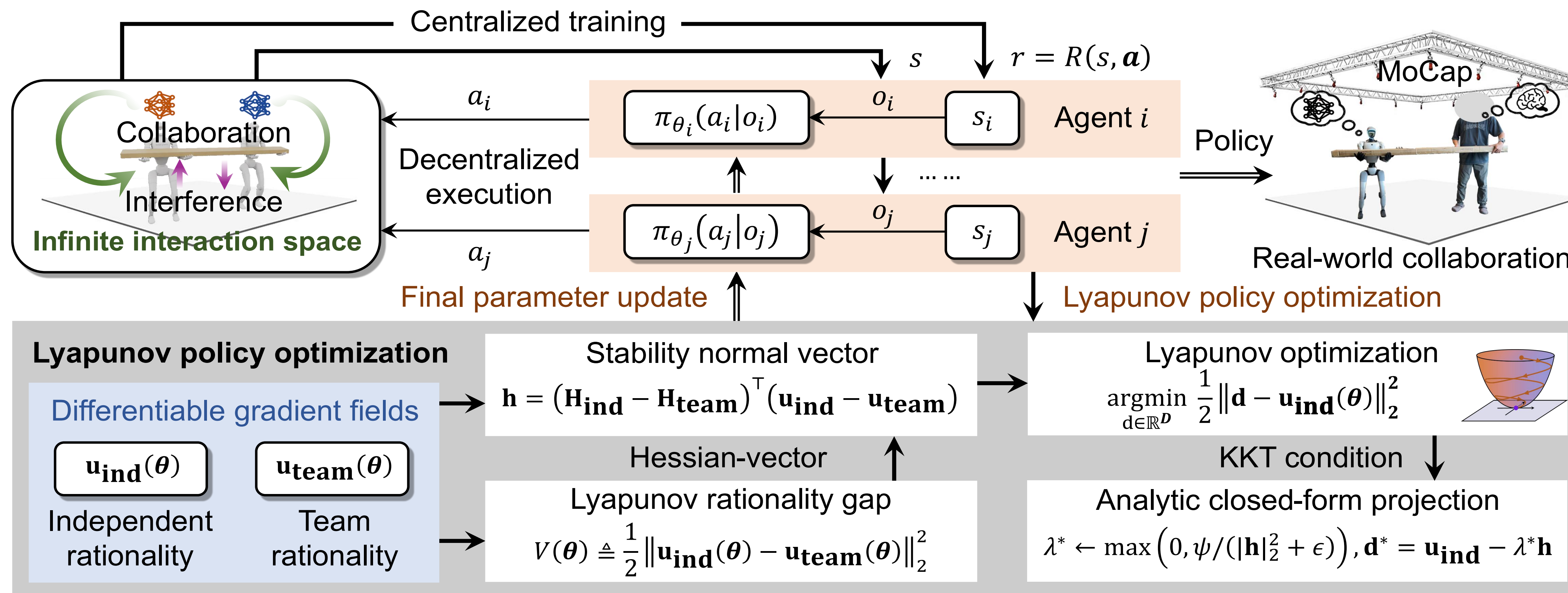
Dr. Hao Zhang, Postdoc Researcher
Carnegie Mellon University
Email: haoz4@andrew.cmu.edu



Project Website

HALO Methodology

We propose **heterogeneous-agent Lyapunov policy optimization (HALO)**, a framework that stabilizes decentralized MARL by enforcing Lyapunov-based contraction in policy-parameter space. Unlike Lyapunov-based safe RL, which targets state/trajectory constraints in constrained Markov decision processes, HALO uses Lyapunov certification to stabilize decentralized policy learning.



RG & Lyapunov Objective

The Rationality Gap $V(\theta)$ – Defined via a Lyapunov function measuring the mismatch:

$$V(\theta) \triangleq \frac{1}{2} \|\mathbf{u}_{\text{ind}}(\theta) - \mathbf{u}_{\text{team}}(\theta)\|_2^2$$

Objective – Enforce a Lyapunov dissipation constraint for asymptotic contraction:

$$\langle \nabla_{\theta} V, \mathbf{d} \rangle \leq -\sigma V(\theta), \quad \sigma > 0$$

Scalability via Hessian-Vector Product

The Bottleneck: Computation of $\mathbf{h} = (\mathbf{H}_{\text{ind}} - \mathbf{H}_{\text{team}})^T (\mathbf{u}_{\text{ind}} - \mathbf{u}_{\text{team}})$ requires $\mathcal{O}(D^2)$ Hessian.

HALO Solution: Bypasses matrix materialization via double back-propagation (create_graph=True), executing a cheap Hessian-vector product (HVP).

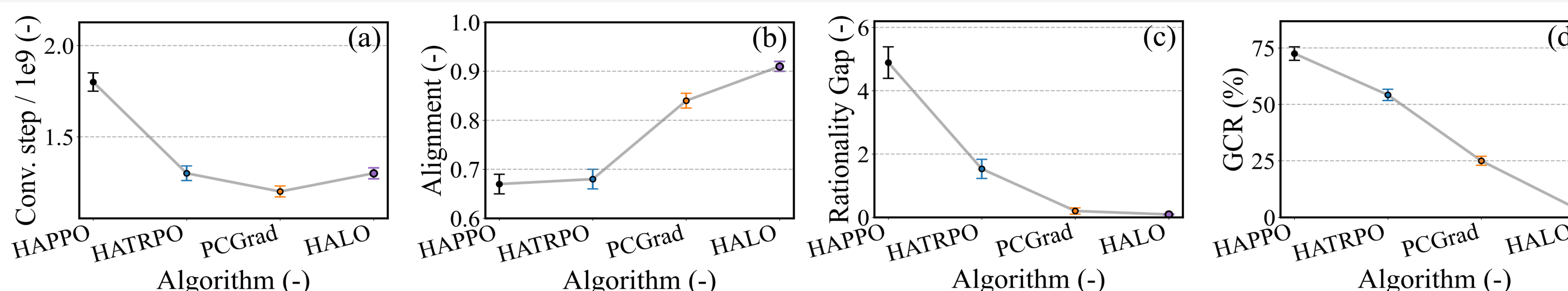
Stability-Constrained Projection

Quadratic Program (QP) – Minimum-norm projection onto stable half-space:

$$\min_{\mathbf{d}} \frac{1}{2} \|\mathbf{d} - \mathbf{u}_{\text{ind}}(\theta_k)\|_2^2 \quad \text{s.t.} \quad \langle \nabla_{\theta} V(\theta_k), \mathbf{d} \rangle \leq -\sigma V(\theta_k)$$

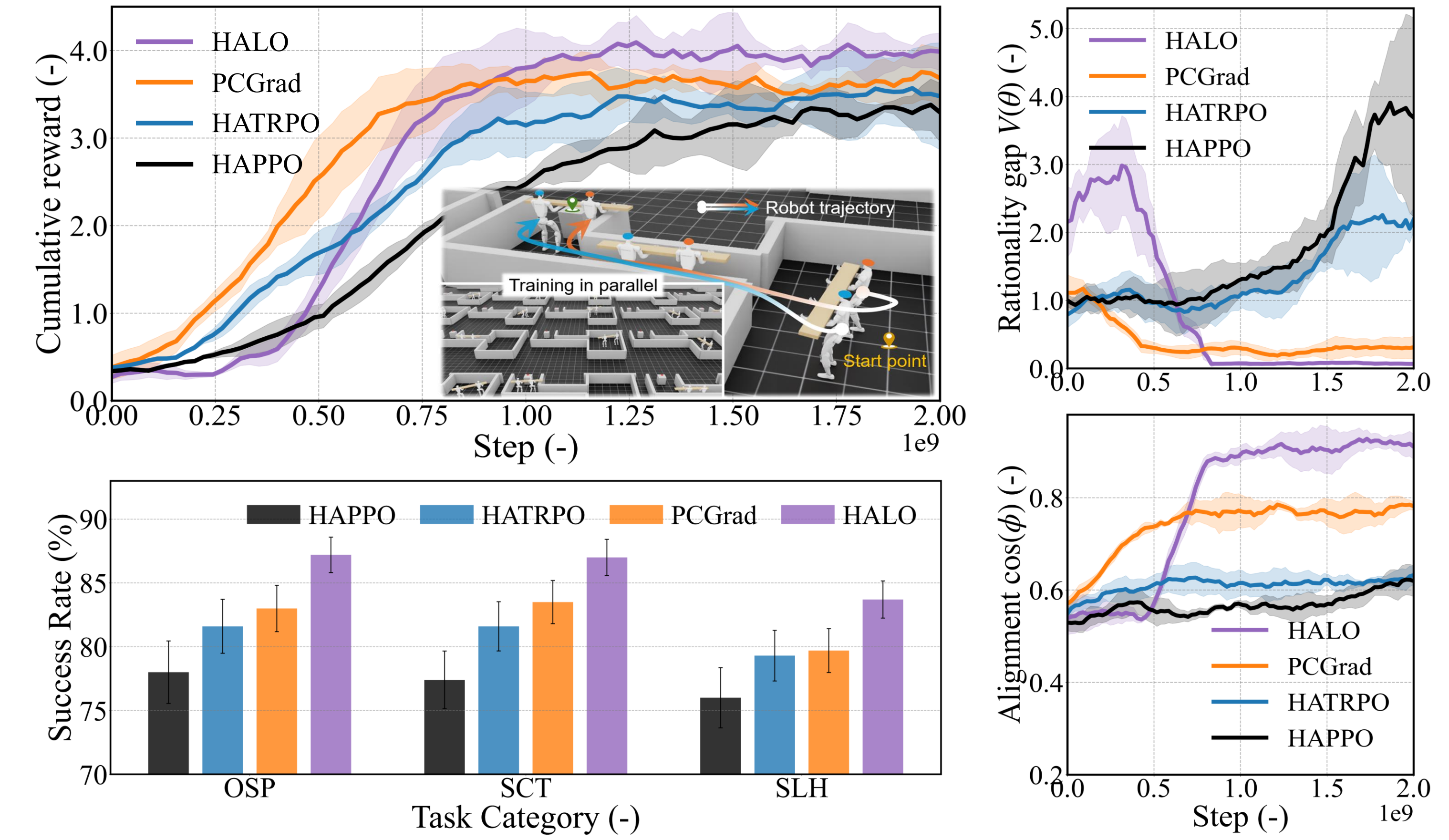
Exact Analytic Solution (KKT) – Here, $\mathbf{h} \triangleq \nabla_{\theta} V(\theta_k)$:

$$\mathbf{d}^* = \mathbf{u}_{\text{ind}} - \max\left(0, \frac{\langle \mathbf{h}, \mathbf{u}_{\text{ind}} \rangle + \sigma V}{\|\mathbf{h}\|_2^2 + \epsilon}\right) \mathbf{h}$$

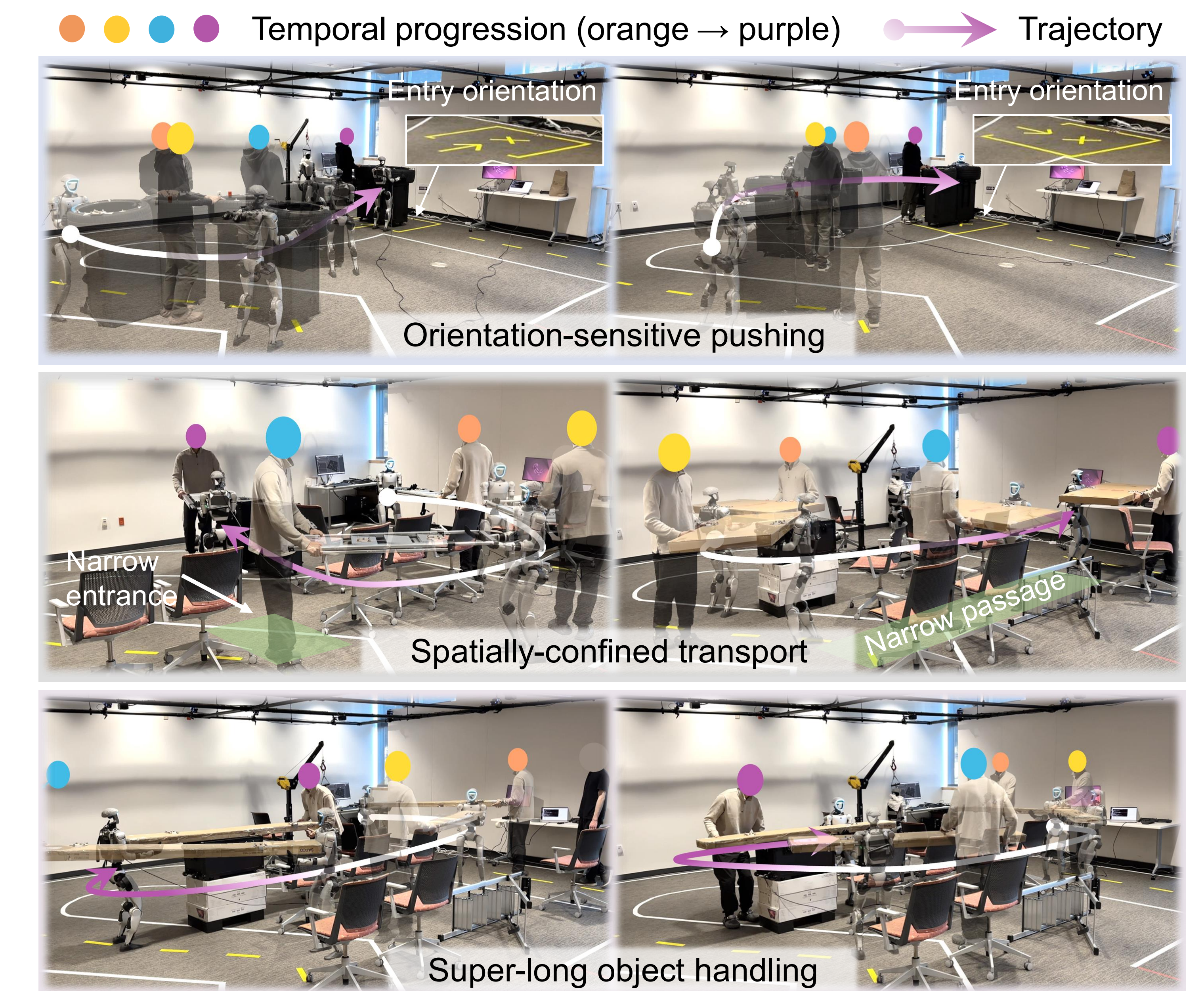


Experiments and Results

Extensive simulations and real-world robot experiments show that this certified stability improves generalization and robustness in collaborative corner case.



Simulation benchmarks and learning dynamics analysis



Sim-to-real deployment across embodied task