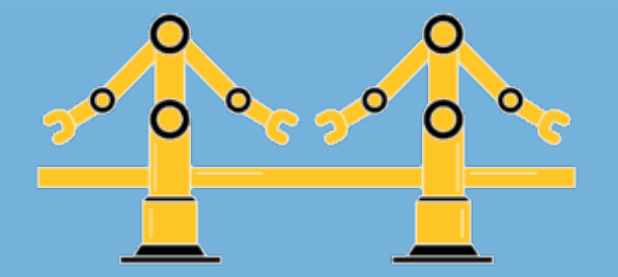




Workers Demo, Robots Learn: Inverse Reinforcement Learning with Cross-Embodiment Videos for Construction Robots



IROS 2026
PITTSBURGH



Lei Huang, Tong Ren, and Zhengbo Zou

Department of Civil Engineering and Engineering Mechanics, Columbia University

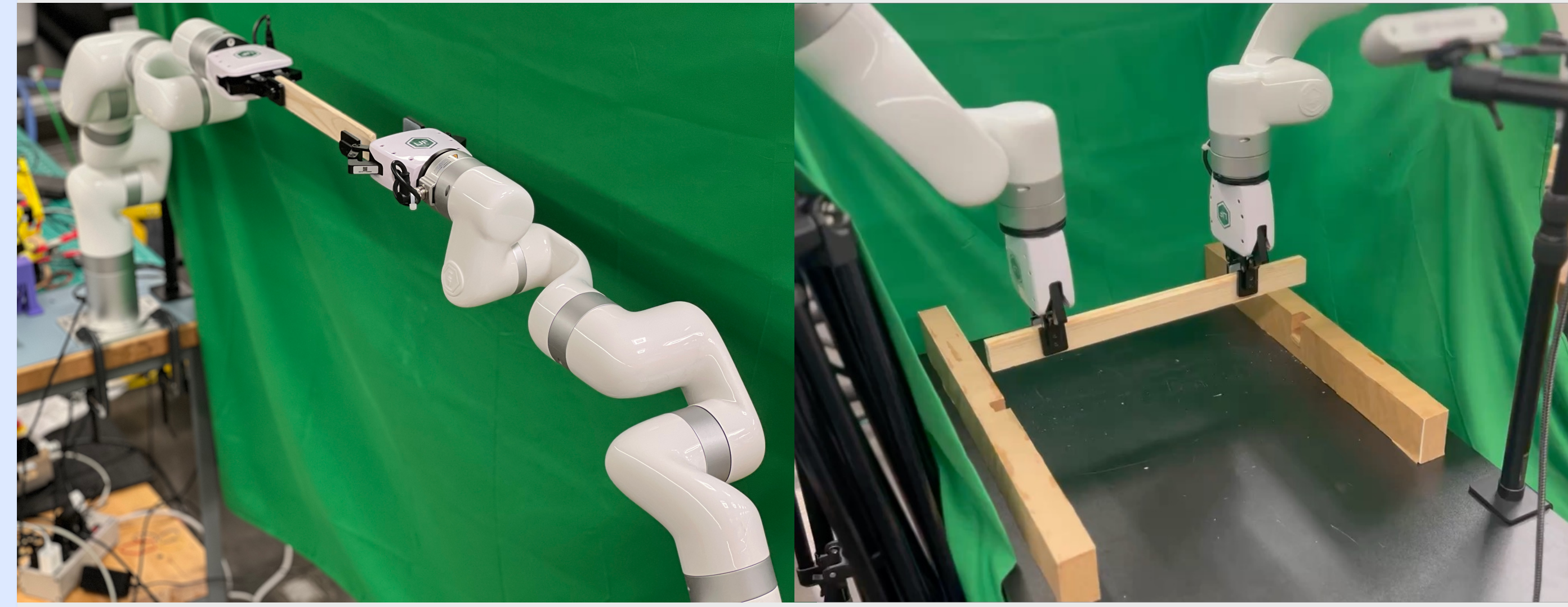
5th Workshop on Future of Construction

Motivation

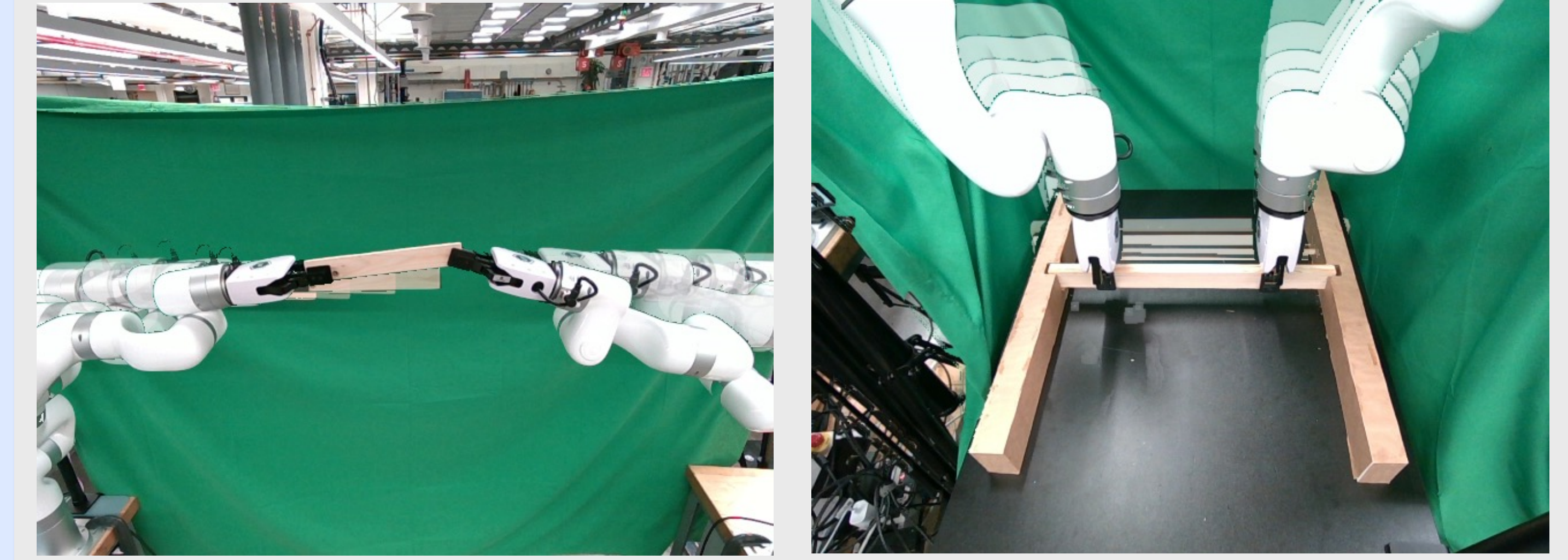
Construction robots obtain higher autonomy by learning policies via **Reinforcement Learning (RL)**, but

- 1) Each policy requires a proper *dense reward function*
- 2) Design of such reward function requires *iterative tuning*
- 3) *Handcrafting reward is not scalable* to the large repertoire of construction tasks

Tested Tasks

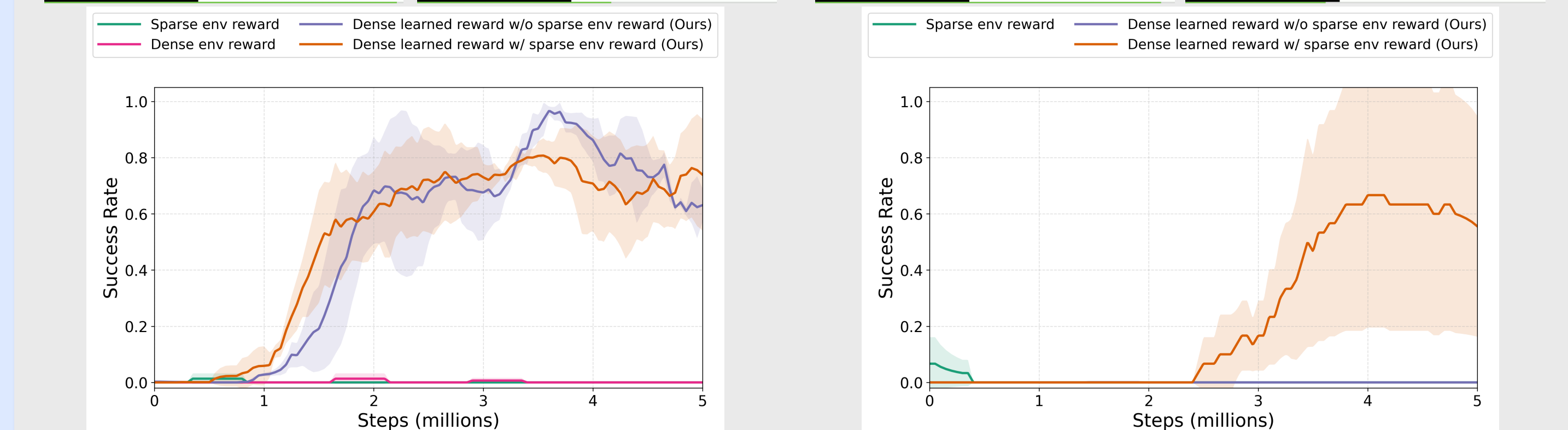
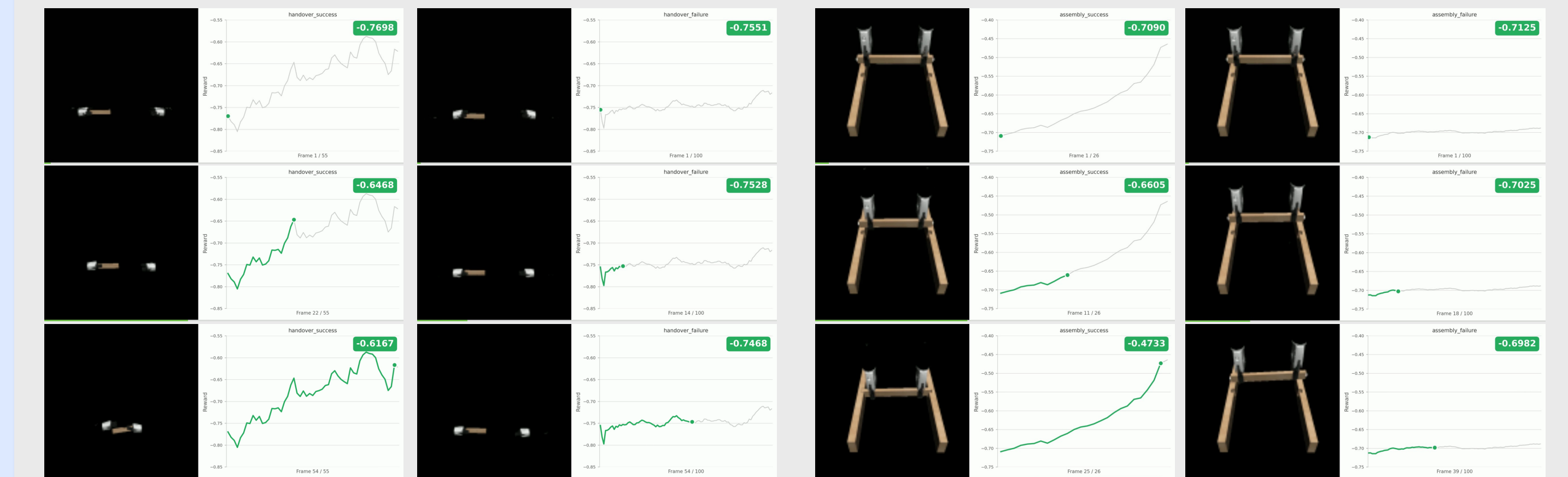


Experiments



Handover Reward: Success vs Failure

Assembly Reward: Success vs Failure

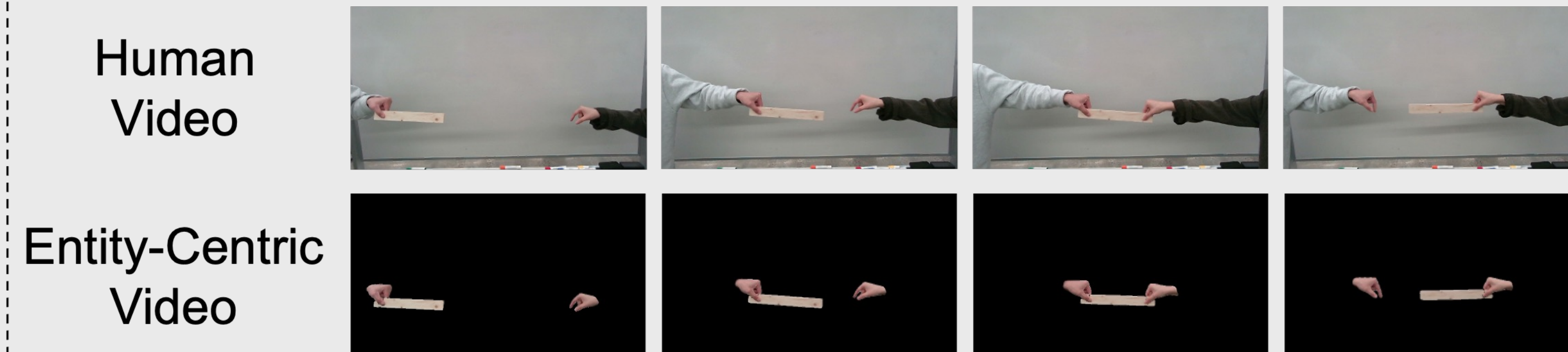


Reward Setting	Success Rate (%)
Sparse Env. Reward	0
Dense Env. Reward	0
Learned Reward	68
Learned Reward w/ Sparse Env. Reward	72

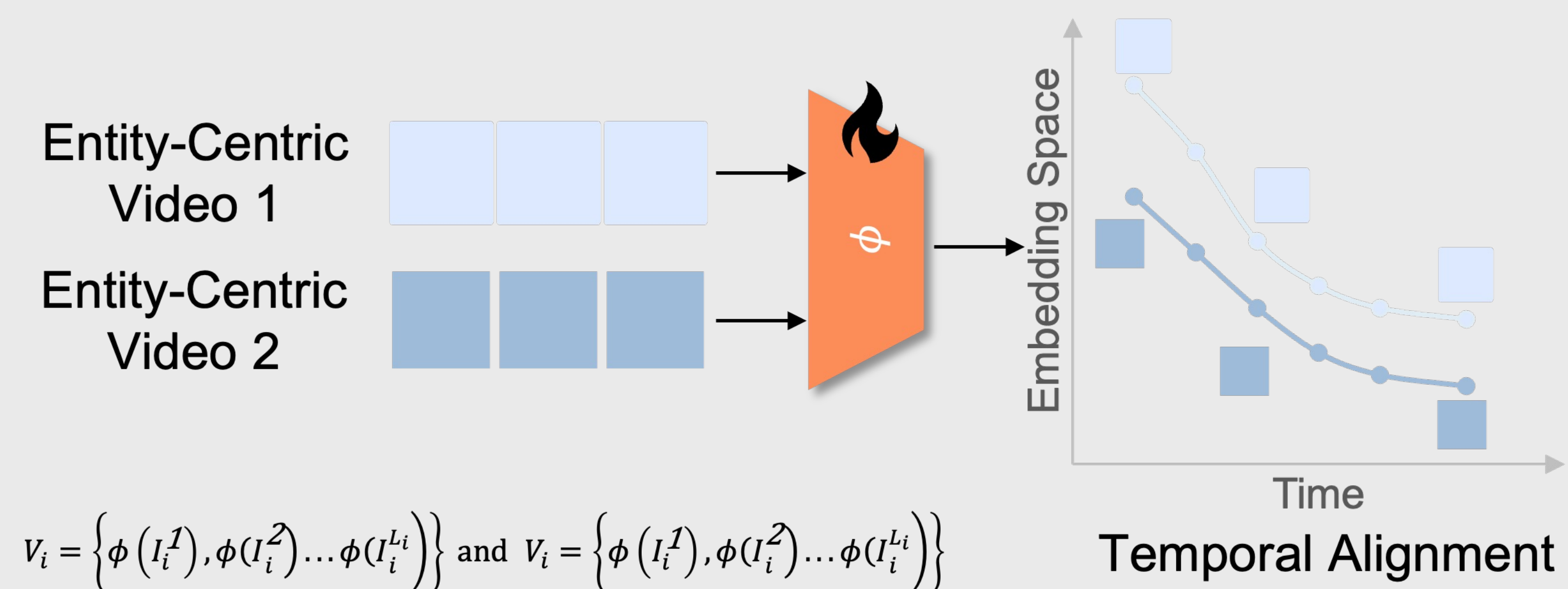
Reward Setting	Success Rate (%)
Sparse Env. Reward	0
Dense Env. Reward	0
Learned Reward	64
Learned Reward w/ Sparse Env. Reward	64

Inverse Reinforcement Learning with Human Videos

1. Object-Centric Segmentation of Human Videos



2. Reward Learning from Human Videos

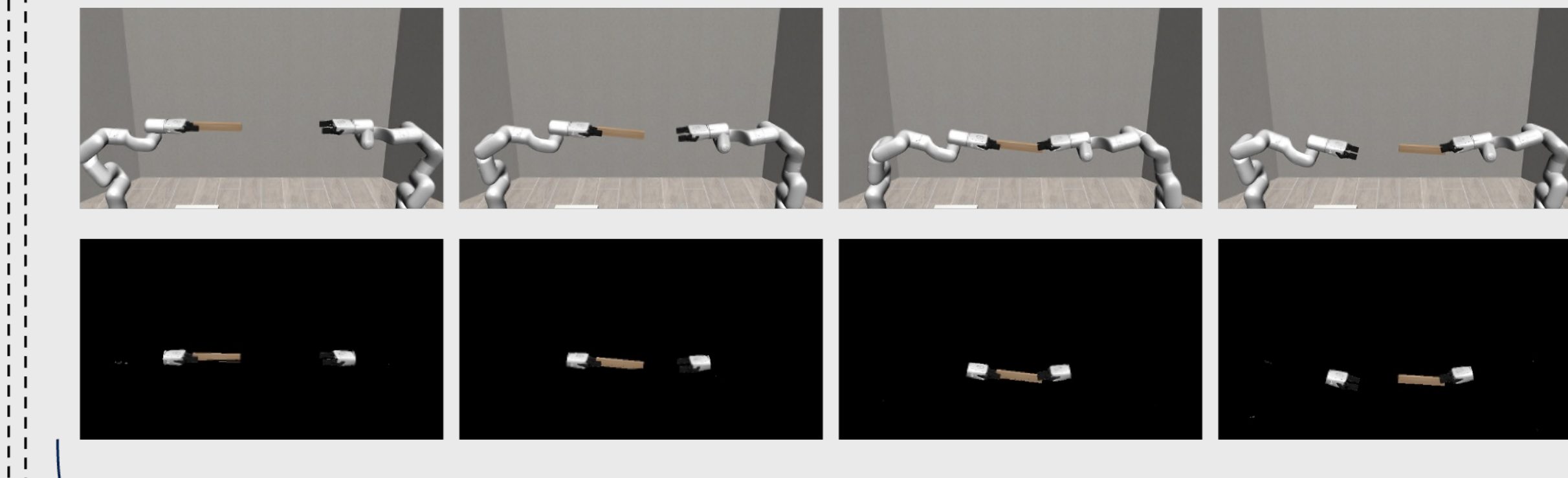


$$V_i = \left\{ \phi(I_i^1), \phi(I_i^2) \dots \phi(I_i^{L_i}) \right\} \text{ and } V_j = \left\{ \phi(I_j^1), \phi(I_j^2) \dots \phi(I_j^{L_j}) \right\}$$

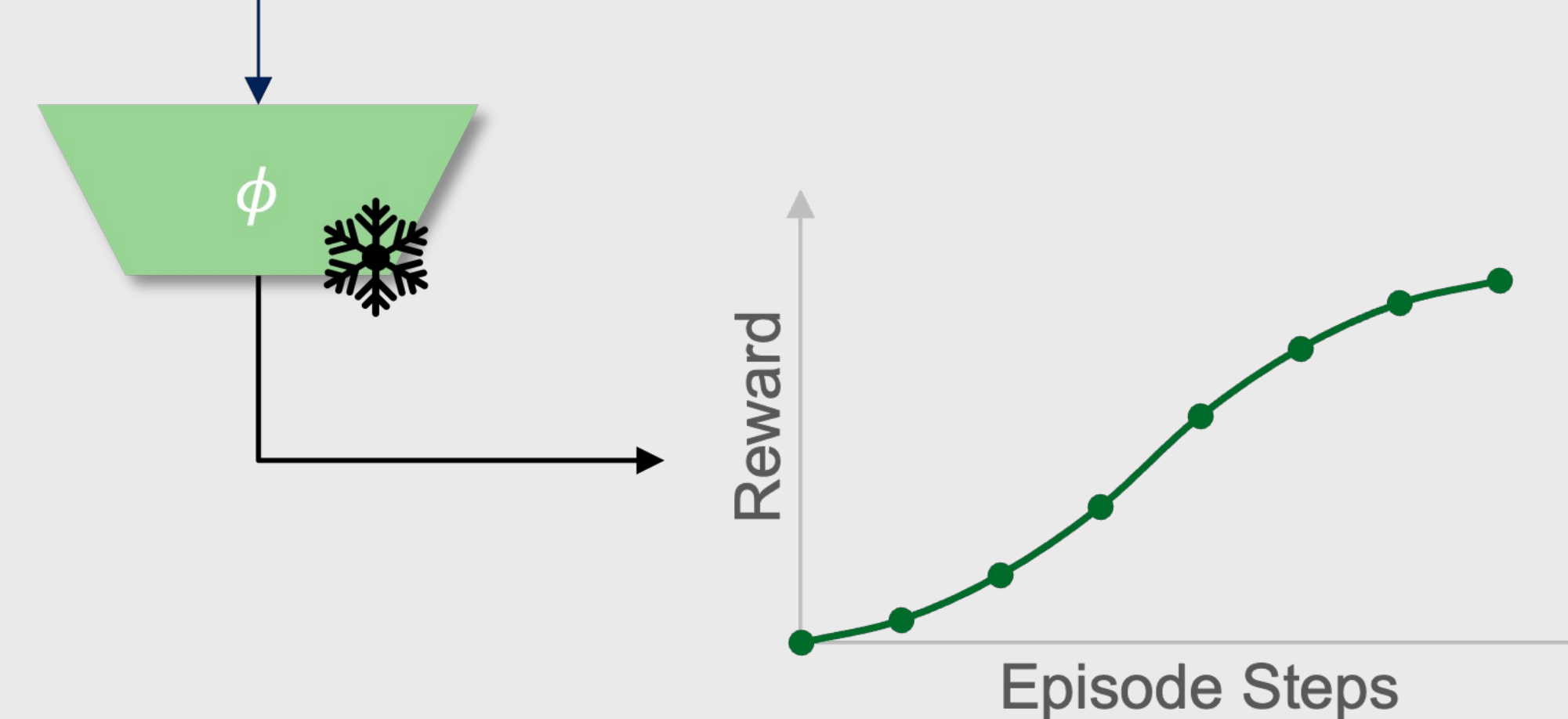
$$\tilde{V}_{ij}^t = \sum_k \alpha_k V_j^k \text{ where } \alpha_k = \frac{e^{-\|v_i^t - v_j^k\|^2}}{\sum_k e^{-\|v_i^t - v_j^k\|^2}} \quad \mu_{ij}^t = \sum_k \beta_{ijt}^k \text{ where } \beta_{ijt}^k = \frac{e^{-\|\tilde{v}_{ij}^t - v_j^k\|^2}}{\sum_k e^{-\|\tilde{v}_{ij}^t - v_j^k\|^2}}$$

$$\text{Loss function: } L_{ij}^t = (\mu_{ij}^t - t)^2$$

3. Reinforcement Learning with Learned Reward



Segmented Episode Rollout



$$\text{Goal Embedding: } g = \sum_{i=1}^N \phi(I_i^{L_i}) / N$$

$$\text{Learned Reward Function: } r(s) = 1/k \|\phi(s) - g\|_2^2$$

$$\text{where } k = 1/N \sum_{i=1}^N \|\phi(I_i^1) - g\|_2$$

Conclusions

- Construction workers' operation videos contain implicit reward signaling task progress
- Temporal encoder learns a latent space where distances to goal embedding reflects the progress
- The latent distances are effective reward for training bimanual construction robots via RL