

## Motivation

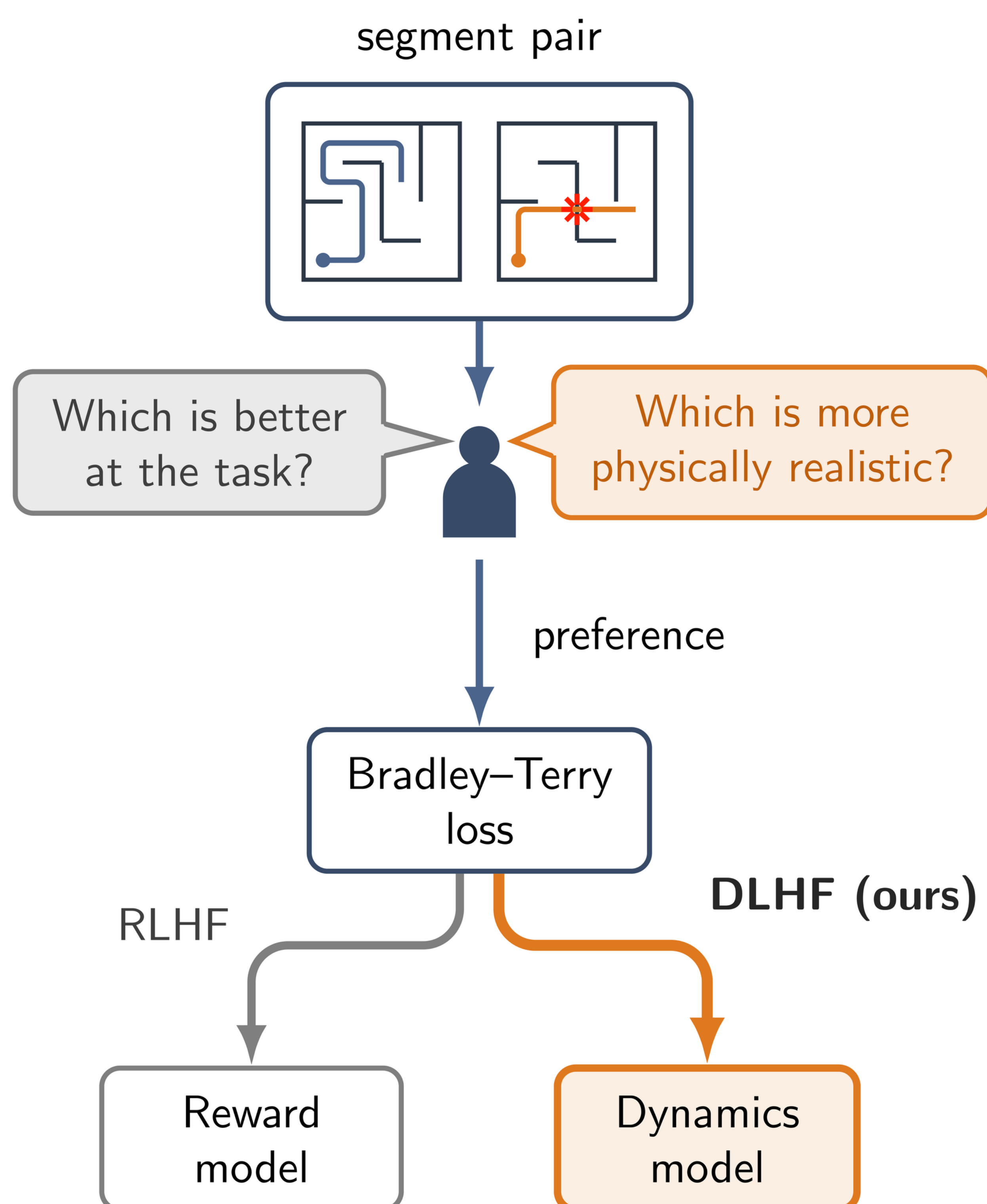
Offline model-based RL suffers **distribution shift** and **model exploitation**. Two typical fixes:

1. Collect more demos → **\$\$\$**, **unsafe**, or **finite**
2. Limit OOD exploration → **Poor generalization**

Can we improve 🌍 without more demos / pessimism?

## Learning 🌍 from Human Feedback

**Insight:** Humans can intuit simple physics!



### RL from Human Feedback (RLHF)

$$L_{DLHF} = \sum \text{CE} \left( y, \text{logistic} \left( \mathbf{R}_\phi(\sigma^0) - \mathbf{R}_\phi(\sigma^1) \right) \right)$$

### Dynamics Learning from Human Feedback (DLHF)

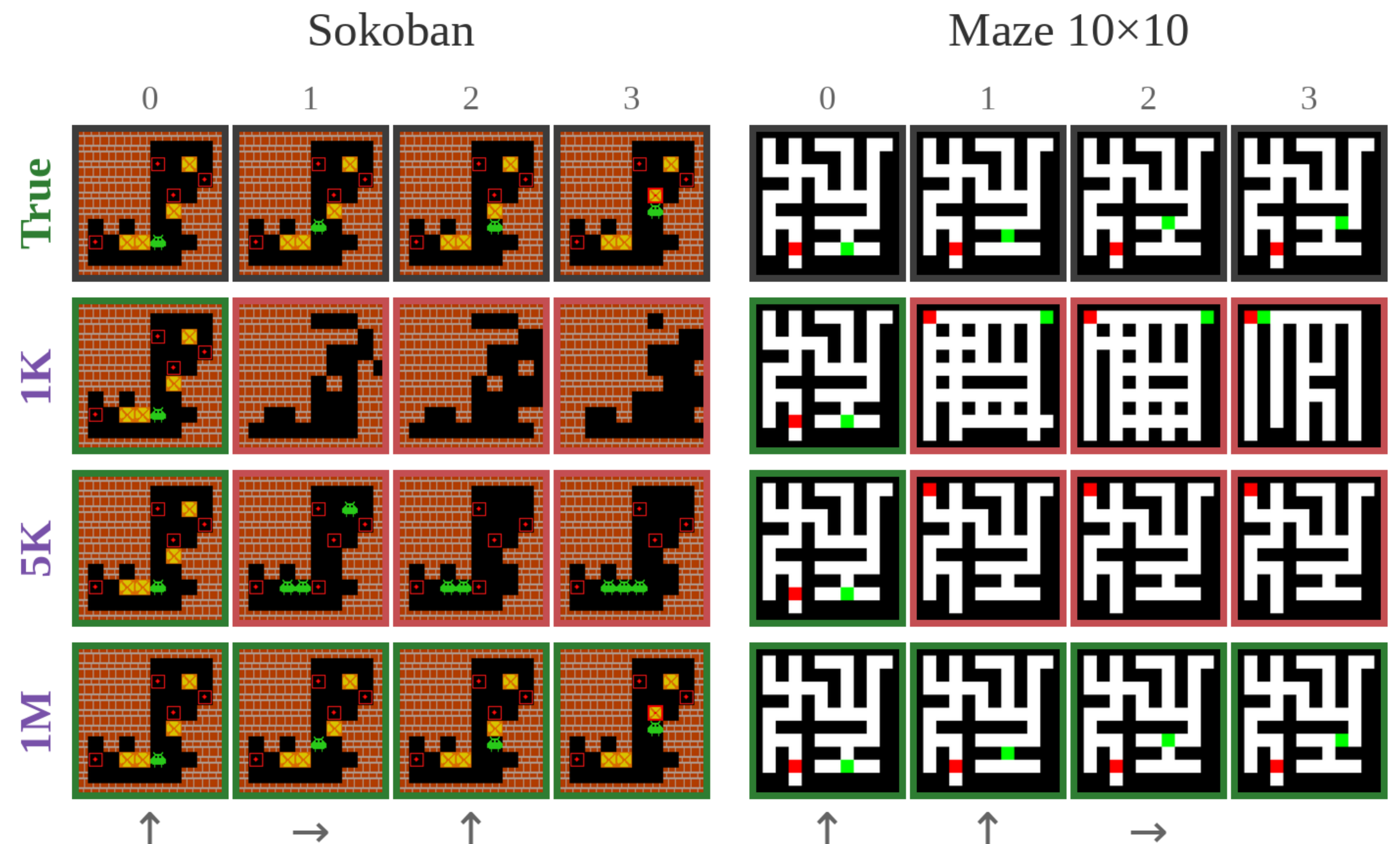
$$L_{DLHF} = \sum \text{CE} \left( y, \text{logistic} \left( \ell_\theta(\sigma^0) - \ell_\phi(\sigma^1) \right) \right)$$

### Algorithm 1 RENEW

**Require:** Pretrained world model  $\hat{T}_\theta$ ,  
preference budget  $N$ , iterations  $I$

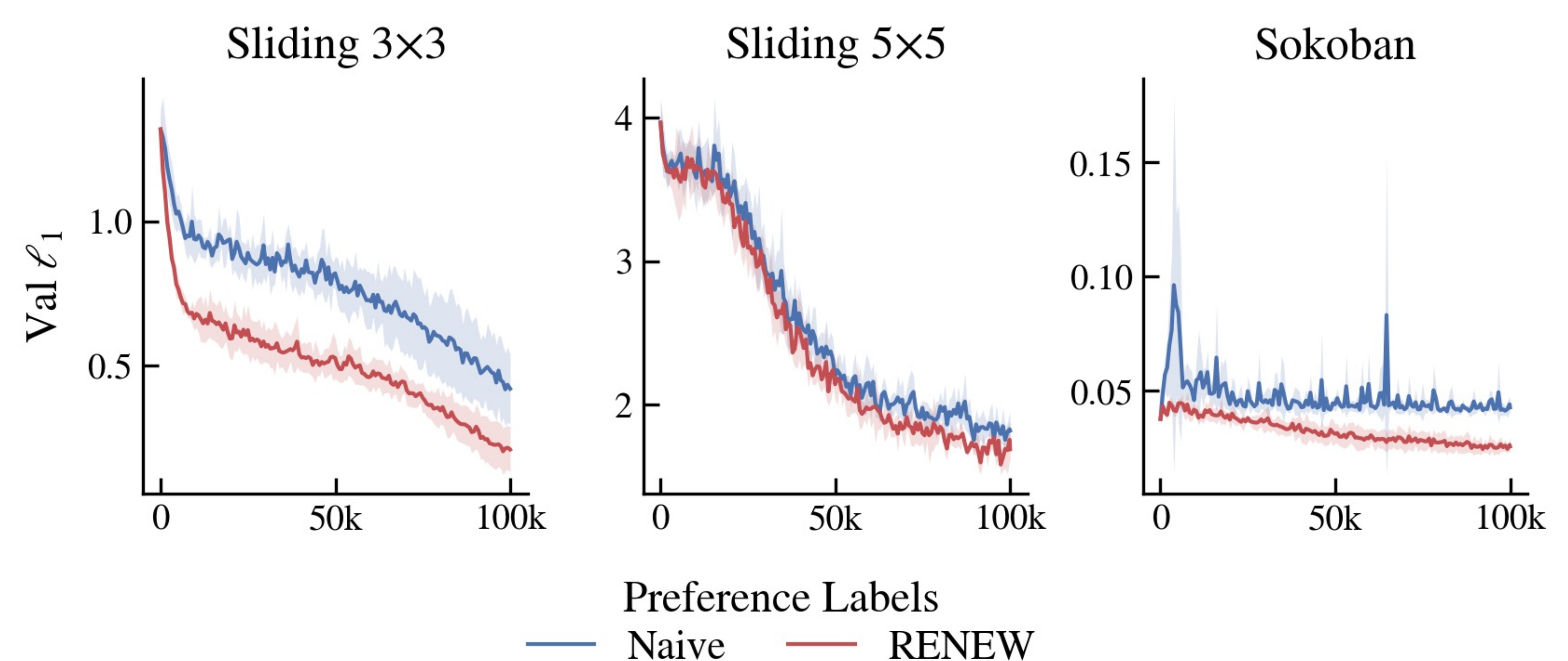
- 1: **for**  $i = 1, \dots, I$  **do**
- 2:   Compute uncertainty  $u$  under  $\hat{T}_\theta$
- 3:   Sample segment pairs  $(\sigma^0, \sigma^1)$  from  $\hat{T}_\theta \propto u$
- 4:   Elicit  $N/I$  preferences  $\mathcal{D}_\succ = \{(\sigma_j^0, \sigma_j^1, y_j)\}_{j=1}^{N/I}$
- 5:   Update  $\theta \leftarrow \arg \min_\theta \mathcal{L}_{DLHF}(\theta; \mathcal{D}_\succ)$
- 6: **end for**
- 7: **return**  $\hat{T}_\theta$

## Q1: Can DLHF learn discrete 🌍?



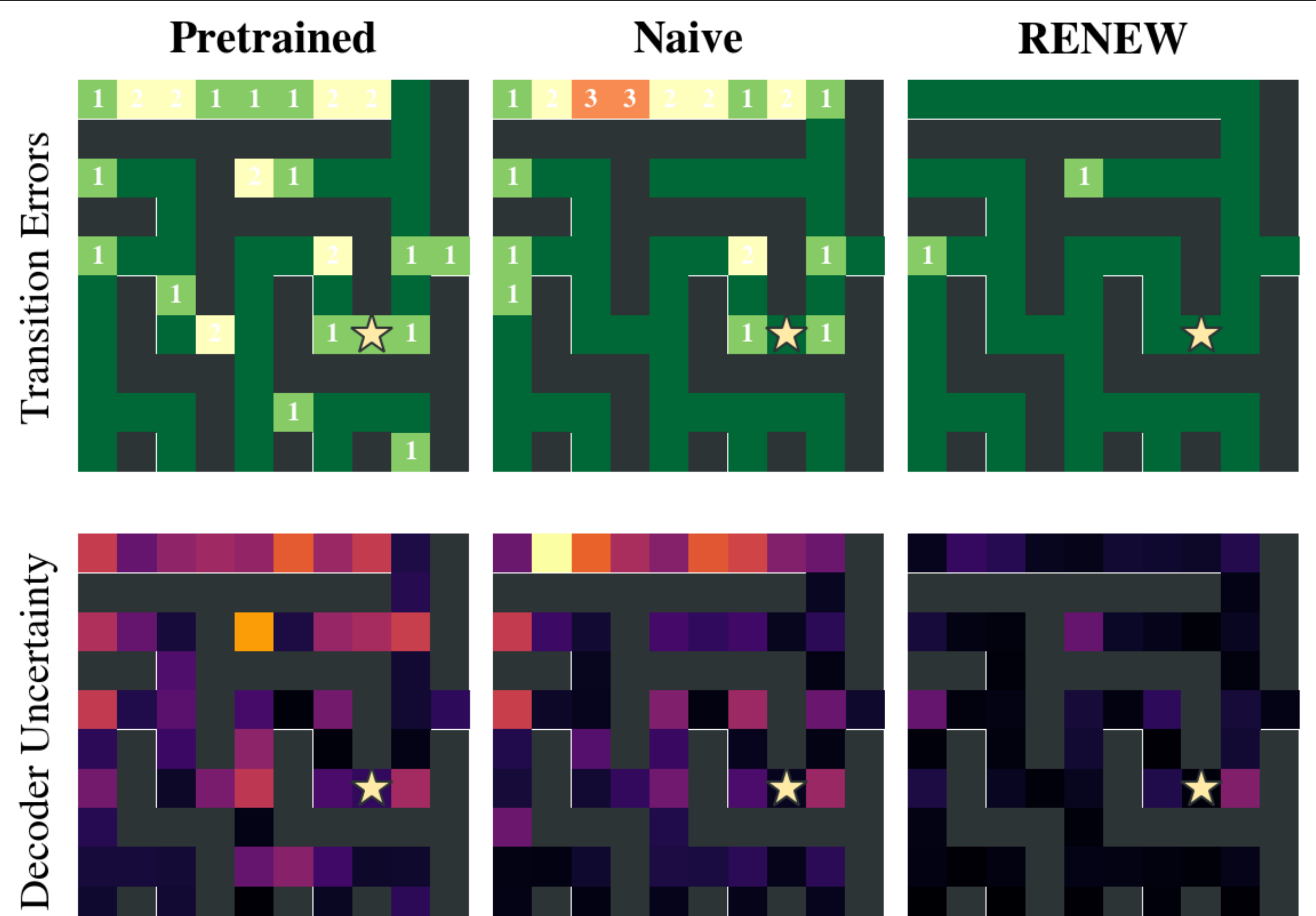
**A1:** Yes, but naïve DLHF is very *sample inefficient*.

## Q2: Can we boost DLHF sample efficiency?



**A2:** Yes, **RENEW** is *more efficient* than naïve DLHF.

## Q3: Can RENEW target model exploitation?



**A3:** Yes,\* **RENEW** *reduces uncertainty & forgetting*.