

# More Than Irrational: Modeling Belief-Biased Agents

Yifan Zhu, Sammie Katt, Samuel Kaski



AAAI-26 / IAAI-26 / EAAI-26



Seemingly irrational decisions can be optimal under biased beliefs from bounded memory  
and we can infer those bounds online to adapt AI assistance

## Problem

Seemingly Irrational Behavior  
in Sequential Decision-Making

### Observation

Humans often take  
**seemingly sub-optimal**  
actions.  
e.g. forget  $\rightarrow$  re-check

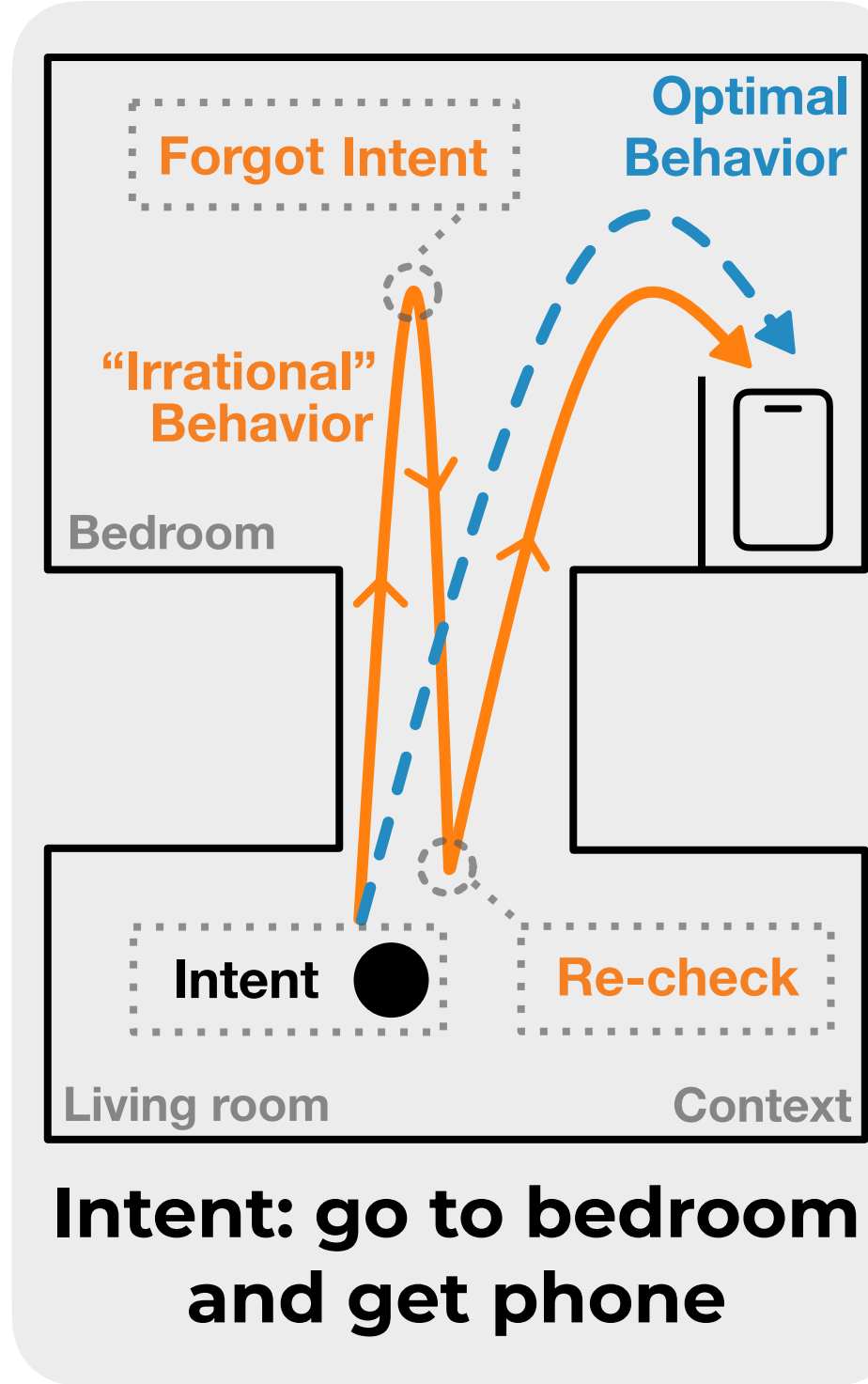
### Challenge

To collaborate with such  
users, the AI needs to

- formalize a cognitively-  
interpretable model
- infer latent factors online

### Insight

Sub-optimal action can be  
**structurally** modeled as the optimal choice under an  
internal biased belief induced by bounded memory.

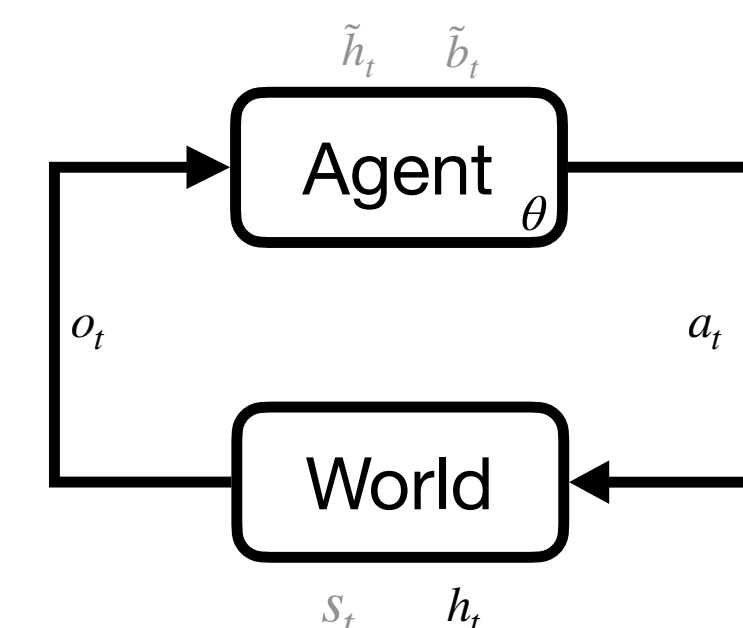


## Model

Computational  
Rationality  
User Modeling

Solved  
By

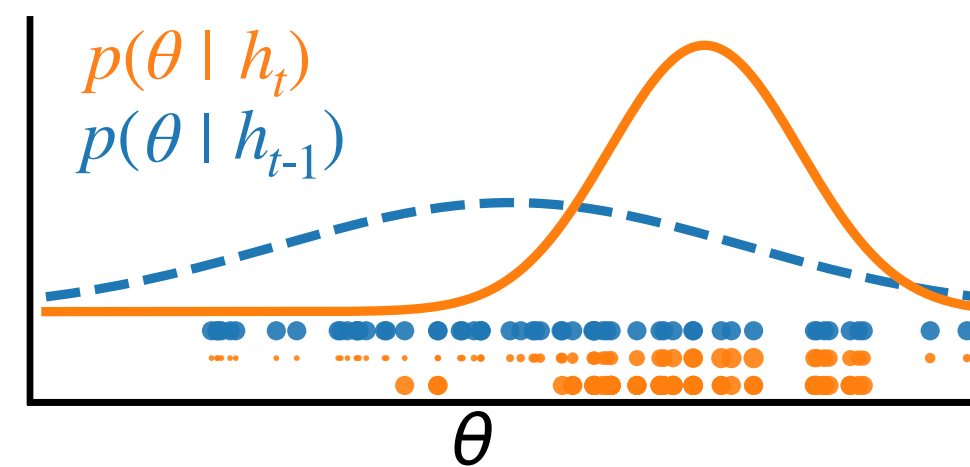
Memory Bound  
 $\rightarrow$  Corrupted Memory  
 $\rightarrow$  Biased Belief  
 $\rightarrow$  Sub-optimal Action



## Infer

Online Inference  
Of  
Bounds and Beliefs

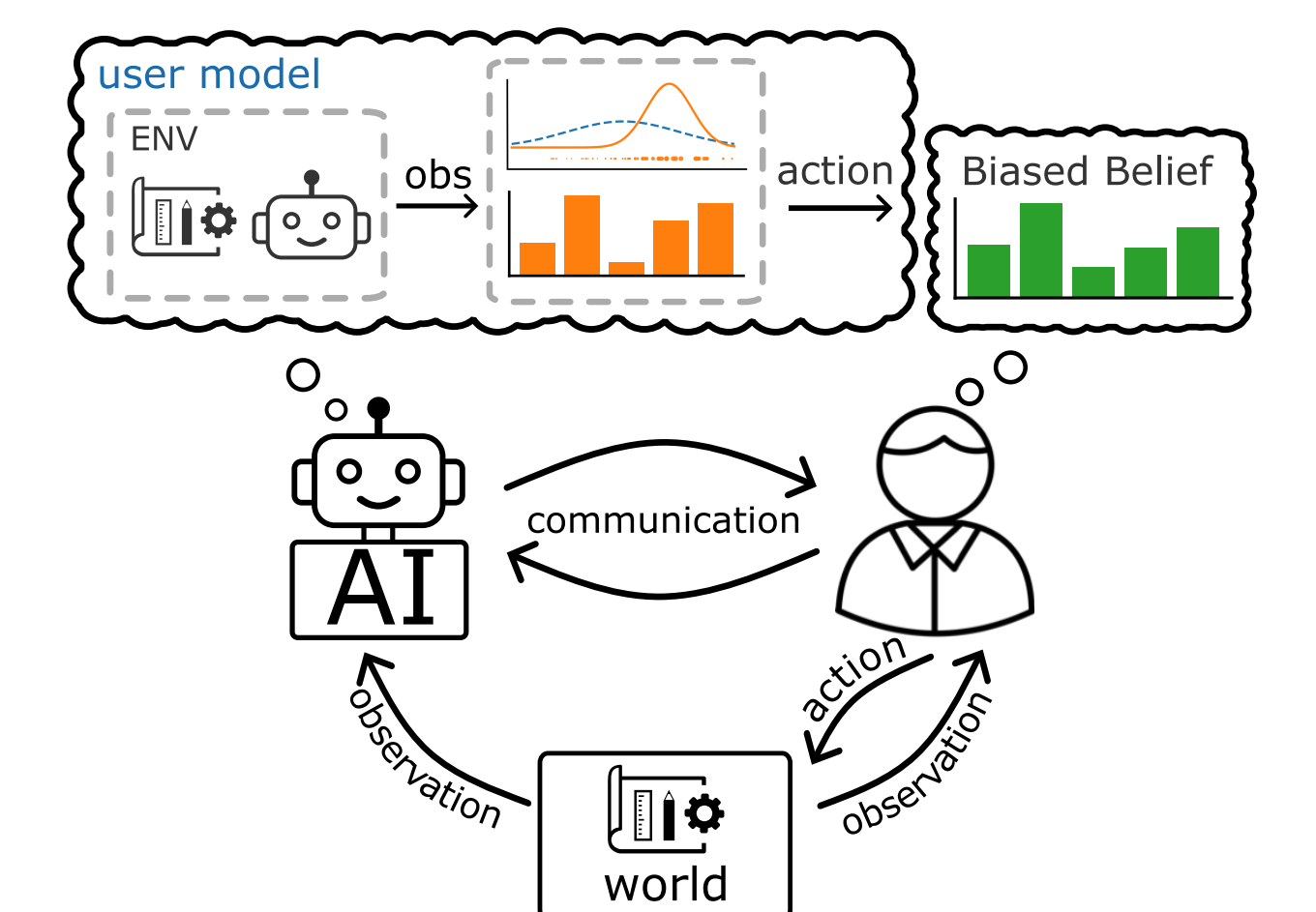
New Data  
(observation, action)  
 $\rightarrow$  Nested Particle  
Filtering  
 $\rightarrow$  Posterior



## Assist

Bound-Belief-Aware  
Adaptive Assistance

Posterior  $\rightarrow$  Cost-sensitive  
Intervention Policy  
 $\rightarrow$  Assistance



## Methodology

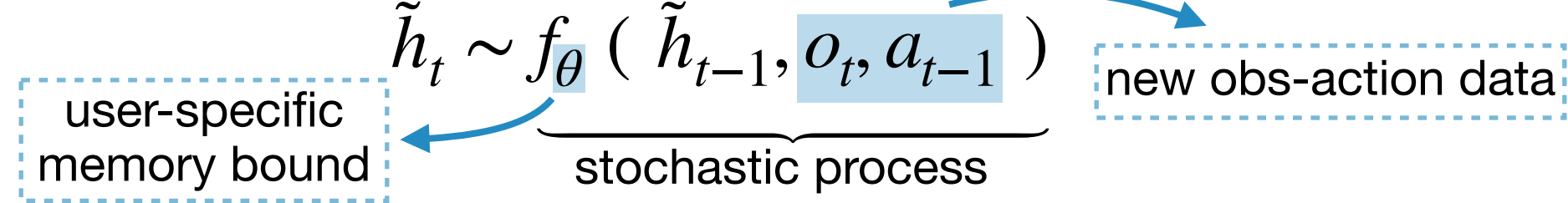
### Computational Rationality User Modeling

Action Likelihood for Inference

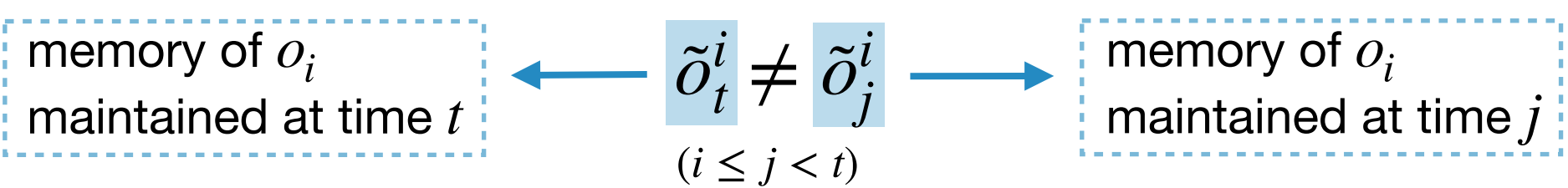
CR-POMDP:  $M_\theta = (S, A, T, R, \Omega, O, \gamma, f_\theta)$ .

At time  $t$ , the CR agent:

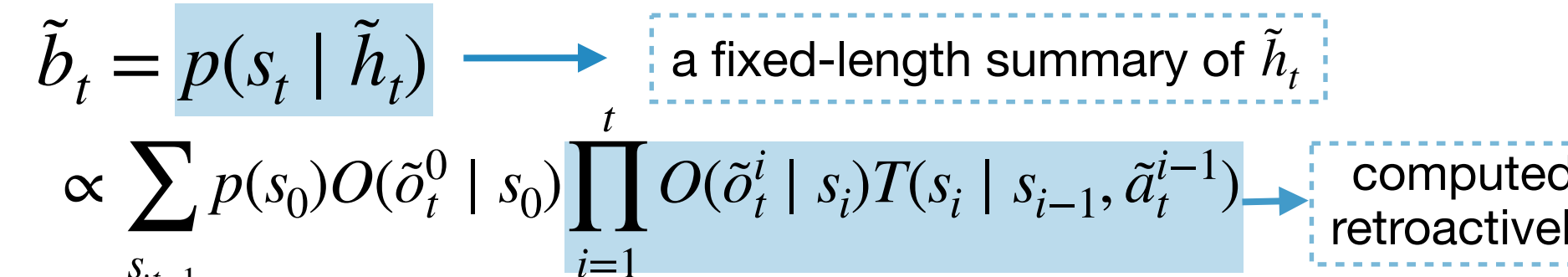
- receives new data, updates **memory**  $\tilde{h}_t = (\tilde{o}_t^i, \tilde{a}_t^{i-1})$ :



- memory of past observations can **evolve**:



- updates **biased belief** from corrupted memory  $\tilde{h}_t$ :



- takes action from the **optimal policy** conditioned on  
the biased belief:  $a_t \sim \pi_*(\cdot | \tilde{b}_t; \theta)$   $\rightarrow$  action likelihood

### Online Inference of Bounds and Beliefs

Nested Particle Filtering

**Inverse Problem:** infer the user's static bound  $\theta$  and latent memory  $\tilde{h}_{t-1}$   
online from observed history  $h_t = (o_t, a_{t-1})$ :

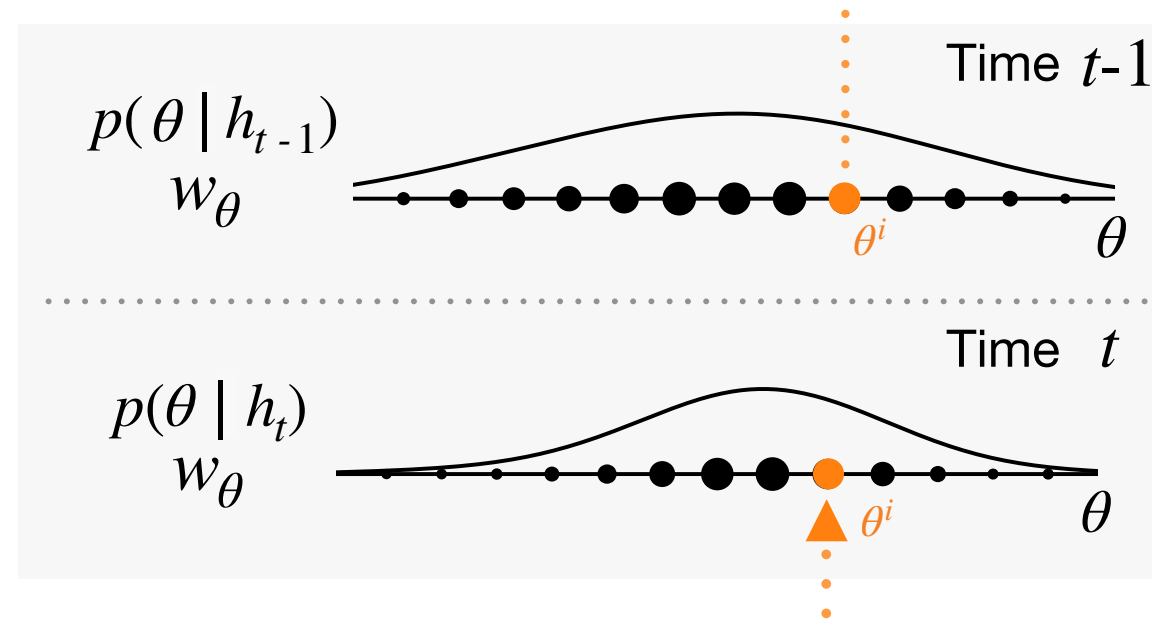
$$p(\tilde{h}_{t-1}, \theta | h_t) \propto p(a_{t-1} | \tilde{h}_{t-1}) \sum_{\tilde{h}_{t-2}} p(\tilde{h}_{t-2}, \theta | h_{t-1}) f_\theta(\tilde{h}_{t-1} | \tilde{h}_{t-2}, o_{t-1}, a_{t-2})$$

action likelihood  $\tilde{h}_{t-2}$  recursive stochastic memory process

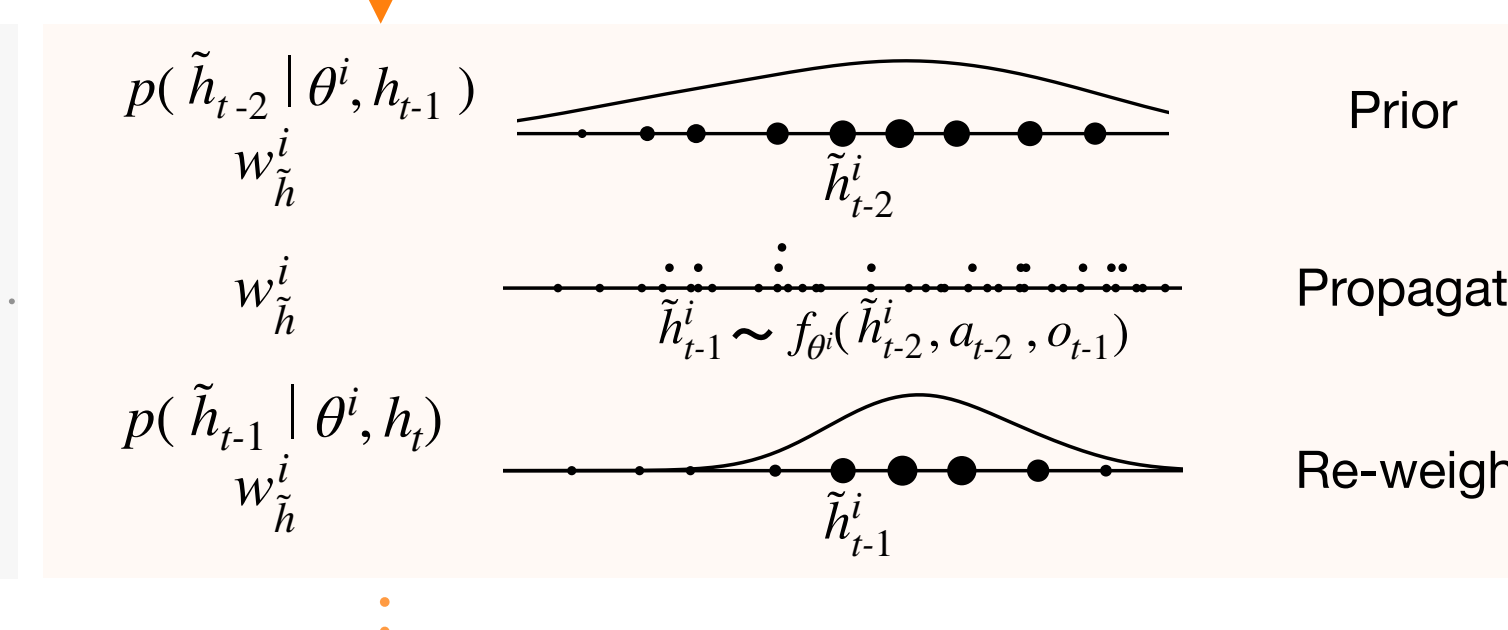
**Challenge:** exact computation is intractable with complexity  $O(|S|^t t!)$ .

**Idea:** approximate  $p(\theta, \tilde{h}_{t-1} | h_t)$  with nested particle filtering; complexity  
comes down to  $O(N_\theta N_{\tilde{h}} t |S|)$ .

Outer particles for  $\theta$



Inner particles for  $\tilde{h}$  under  $\theta$



### Bound-Belief-Aware Adaptive Assistance

Decide when & How to assist  
via inferred bounds and beliefs

**Goal:** learn an assistive policy s.t.

$$\pi_*^{AI} = \arg \max_{\pi^{AI}} \mathbb{E} \left[ \sum_t \gamma^t (r_t - c(a_t^{AI})) \right]$$

task reward  $\downarrow$  intervention cost

**Assistive-POMDP:**

At time  $t$ , the AI:

- observes world and agent:  
 $o_t^{AI} = (o_t, a_{t-1}^H)$
- updates belief via NPF:  
 $b_t^{AI} = p(\theta, \tilde{h}_{t-1} | h_t)$
- proposes intervention from  
{DoNothing, ActionHint, MemoryHint}:  
 $a_t^{AI} \sim \pi_*^{AI}(\cdot | b_t^{AI})$

## Behavioral Validation

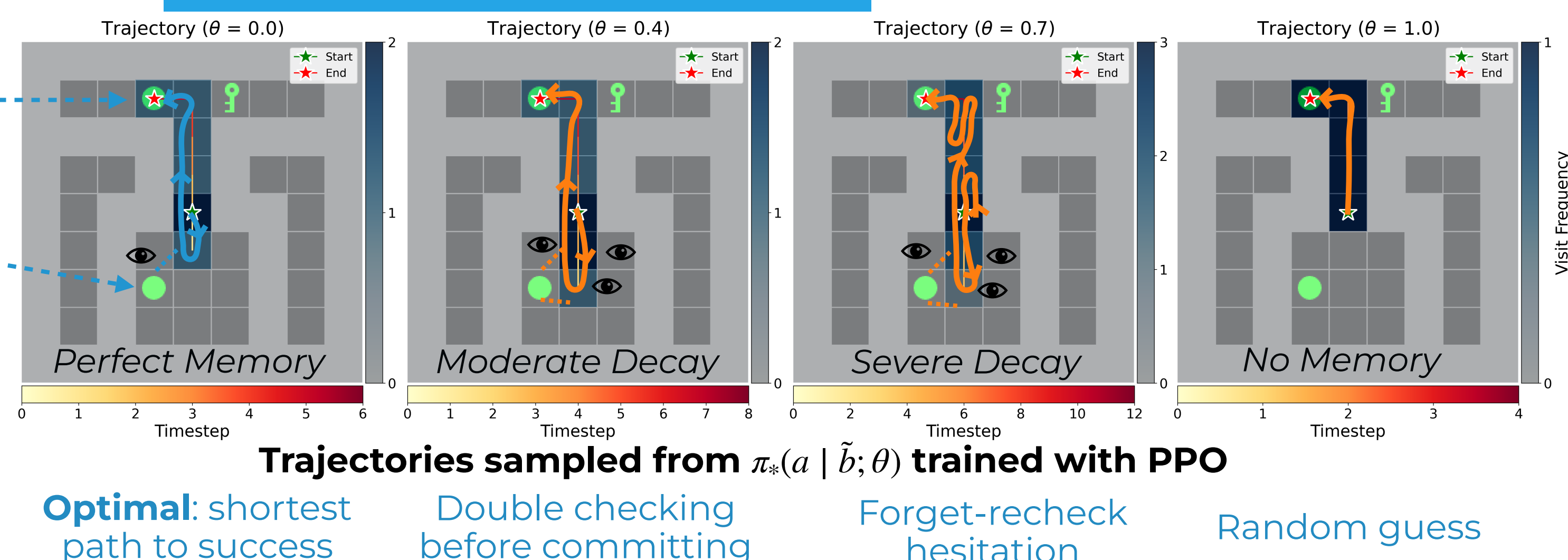
### Setup: T-Maze

**Goal:** reach the correct  
terminal quickly.

The agent must:

- find clue for direction
- remember the clue
- go to the terminal

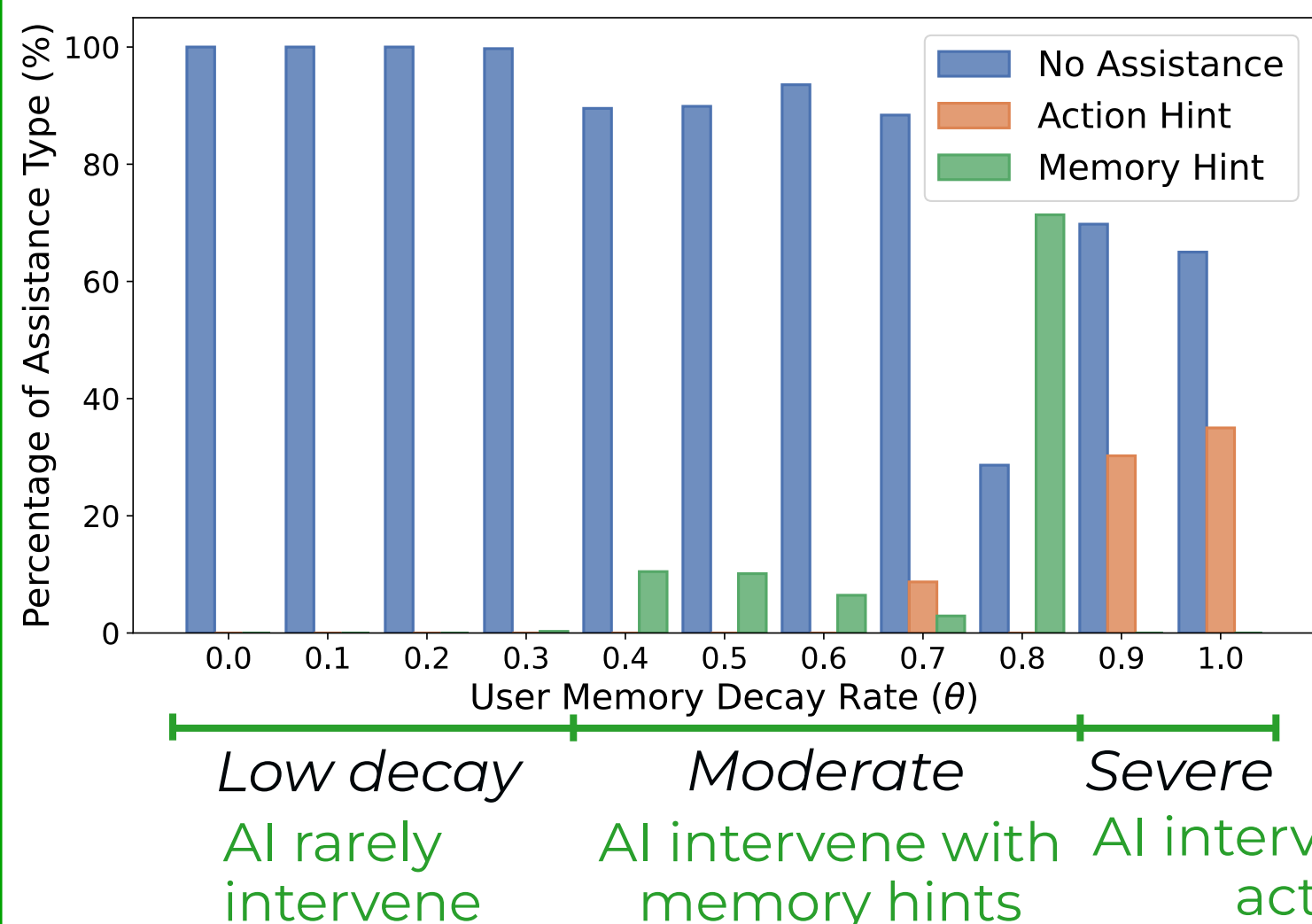
**Memory decay  $f_\theta$ :** a  
stochastic forget process  
with  $\theta$  being decay rate.



## Assistance Validation

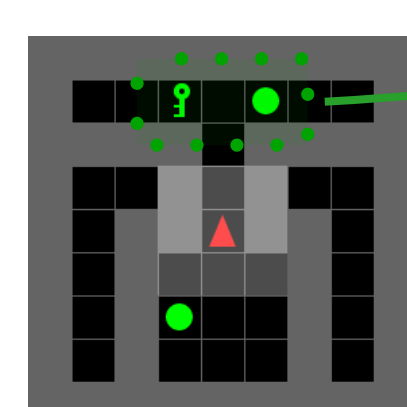
### Type adaptation

Learned Adaptive Assistance Strategy



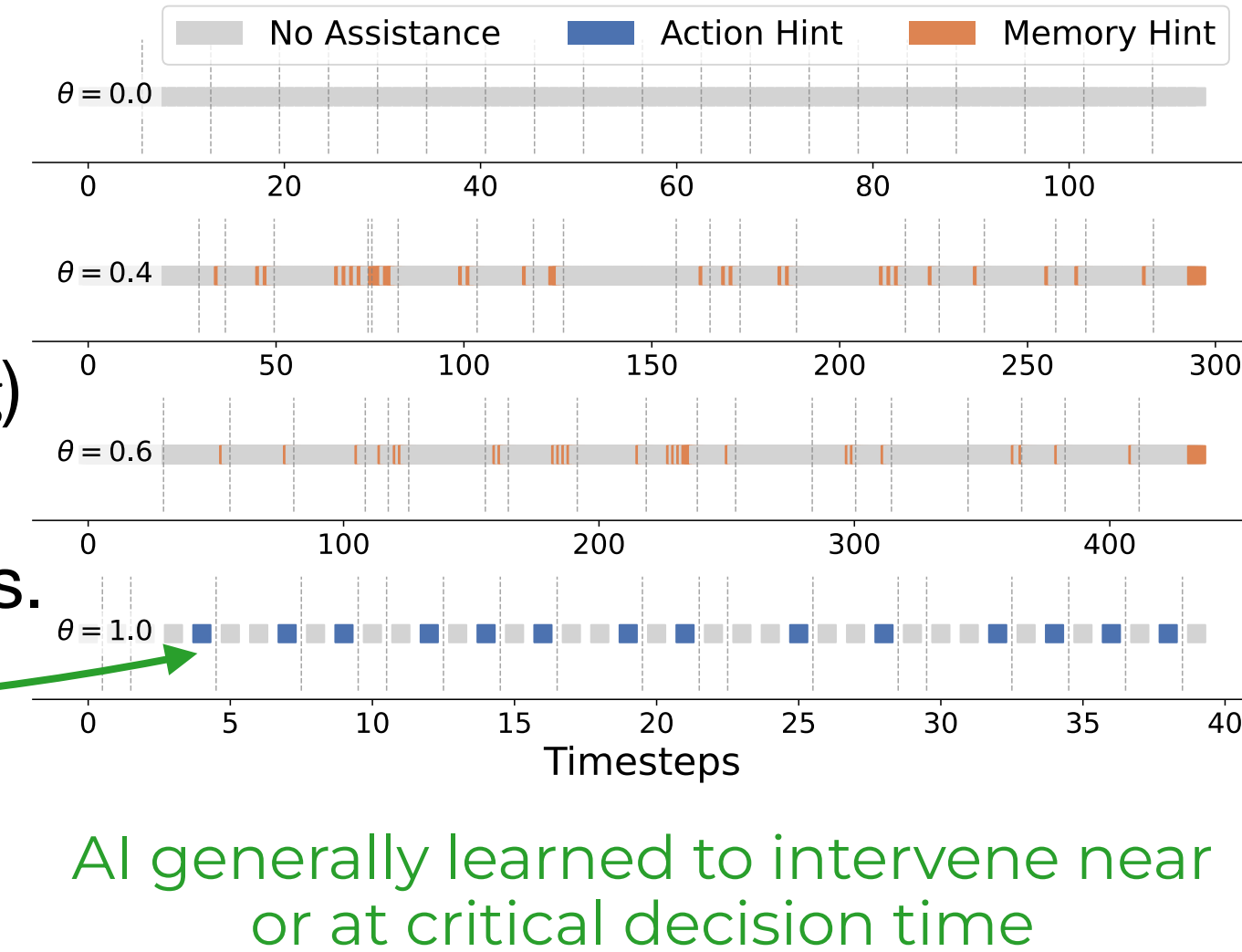
### Setup: T-Maze

- Policy learned with PPO  
using inferred posterior.
- Cost:  $c(\text{ActionHint}) >$   
 $c(\text{MemoryHint}) > c(\text{DoNothing})$
- Critical decision time: last  
step before terminal states.



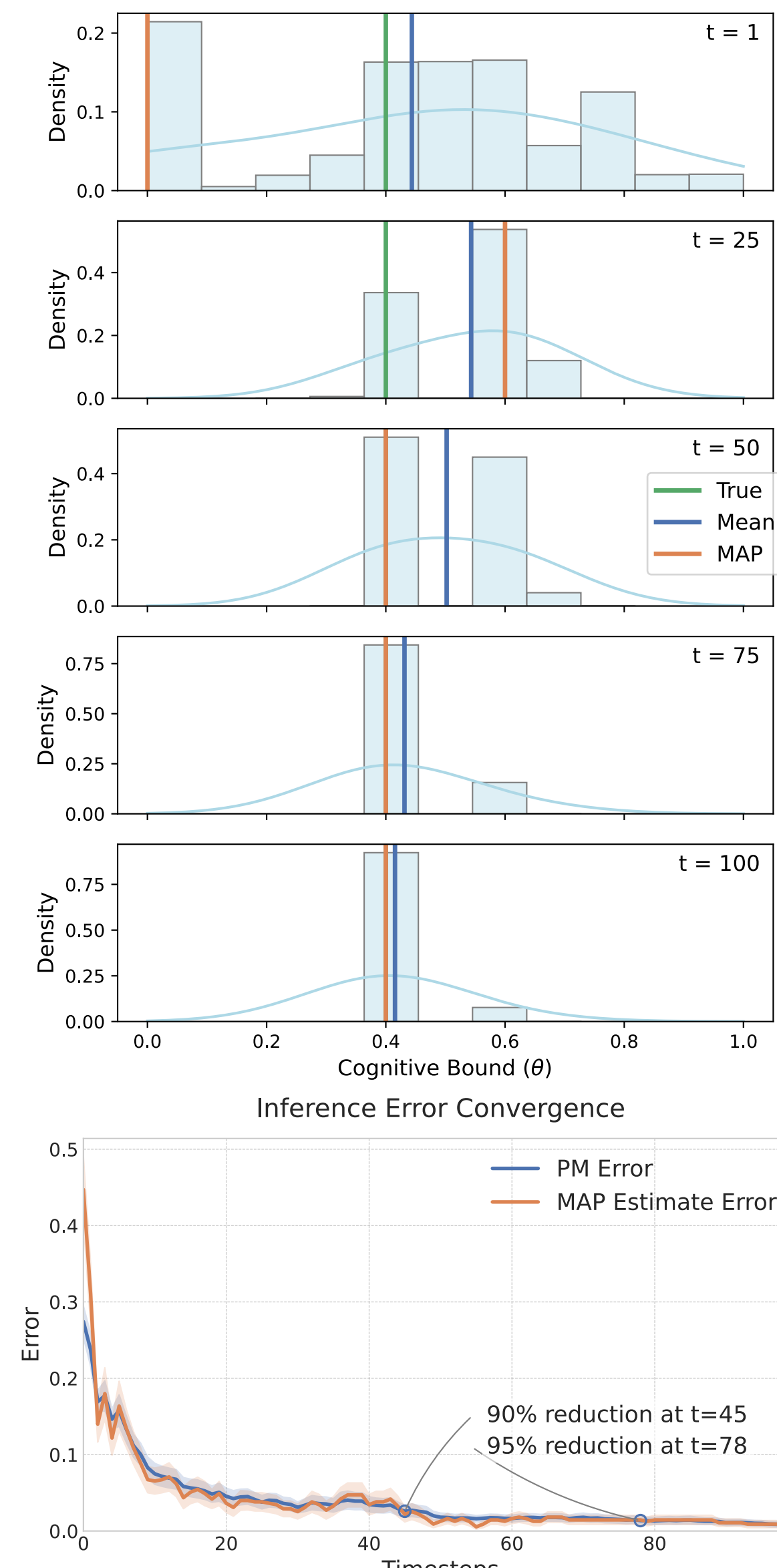
### Timing adaptation

Assistance Sequences for Selective Users



## Inference Validation

Posterior  $P(\theta | h_t)$  Evolution



### Efficacy

NPF quickly  
concentrates  
 $p(\theta | h_t)$  on a small  
set of hypotheses.

### Efficiency

Posterior-mean error  
drops 90% by  $t=45$   
and 95% by  $t=78$ .

Find in the Paper:

- Detailed formulation  
of assistive-POMDP.
- More experiment  
and result details.



Full Paper