



MAX PLANCK INSTITUTE
FOR SOFTWARE SYSTEMS

Robust In-Context Reinforcement Learning Under Reward Poisoning Attacks

Paulius Sasnauskas^{† 1 2} Yiğit Yalın³ Goran Radanović³

[†] This work was done as a part of an internship project at MPI-SWS

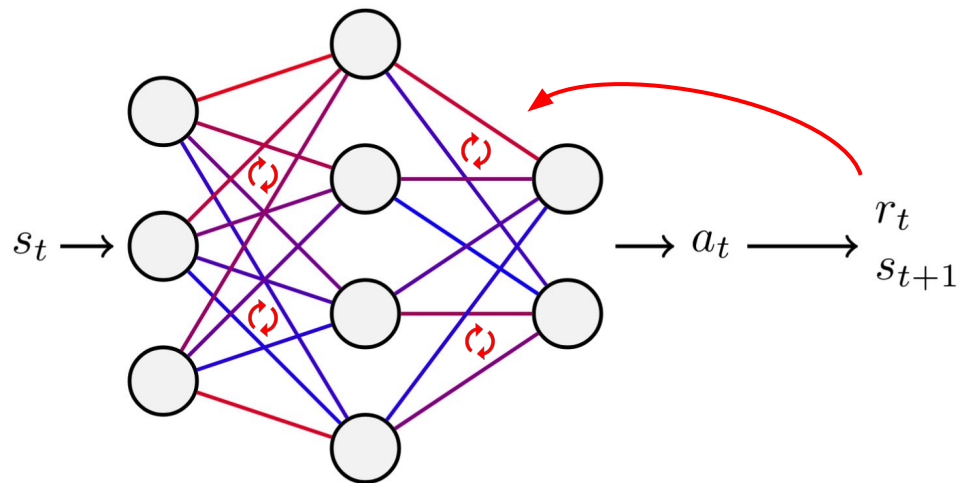
¹ Department of Computing Science, University of Alberta, Edmonton, Canada

² Alberta Machine Intelligence Institute (Amii), Edmonton, Canada

³ MPI-SWS, Saarbrücken, Germany

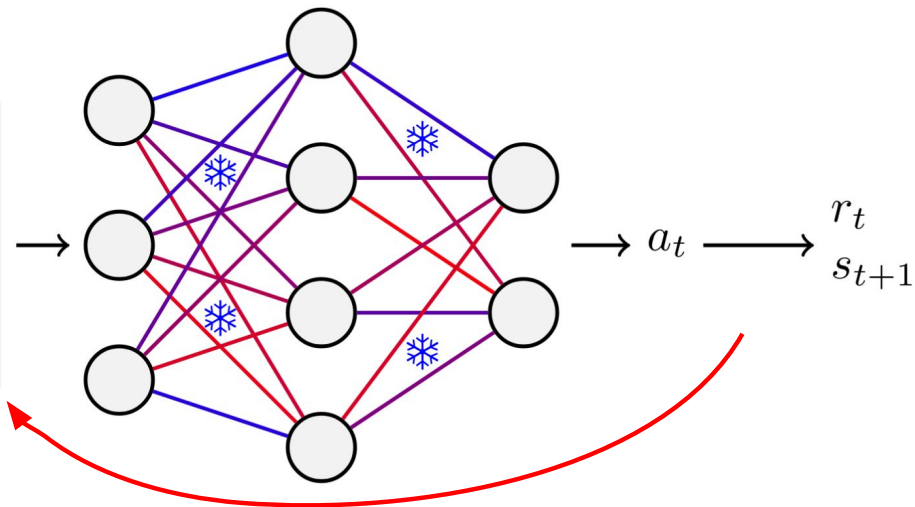
Correspondence to: Paulius Sasnauskas <sasnausk@ualberta.ca>

Single-Task RL:

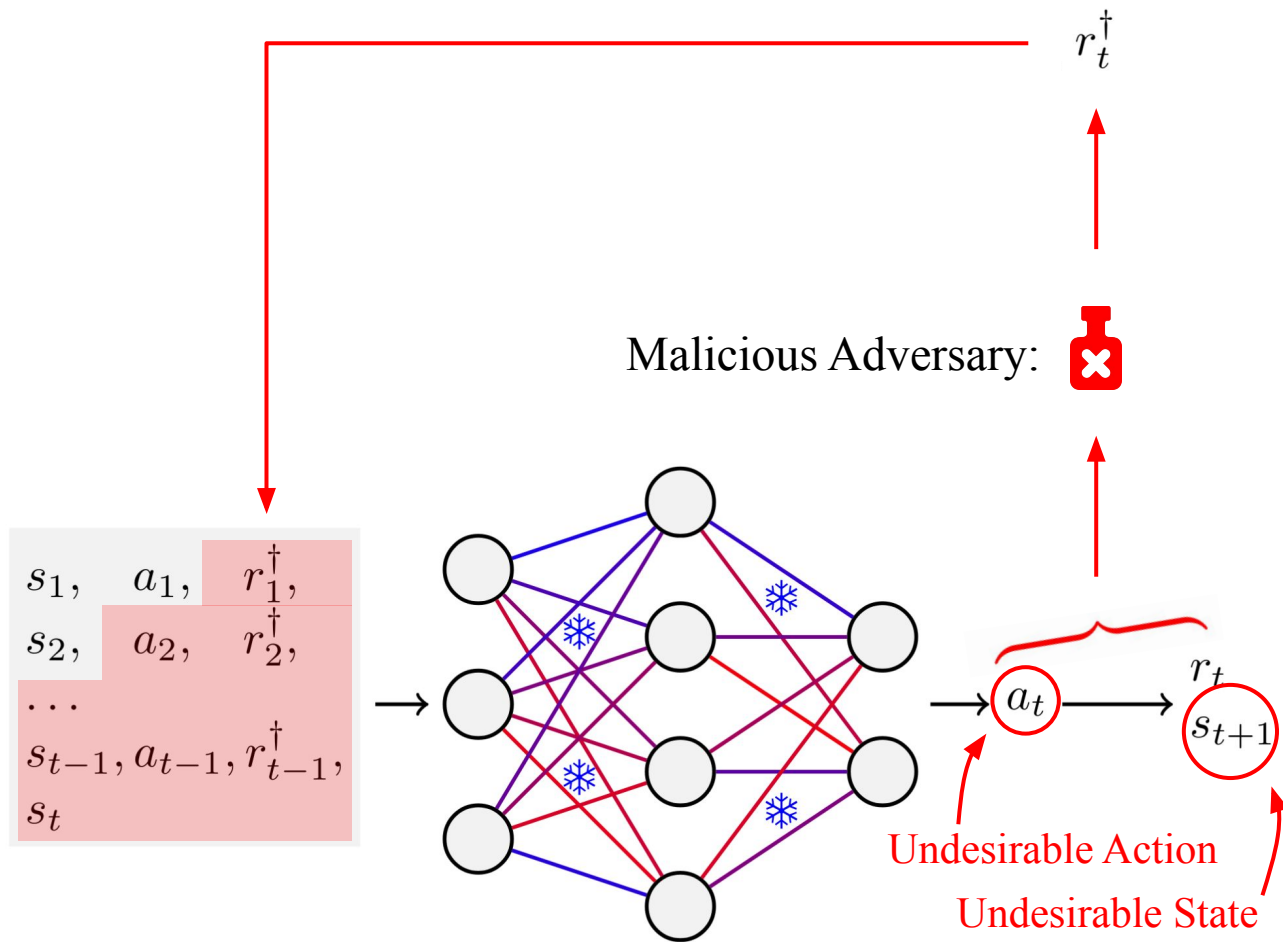


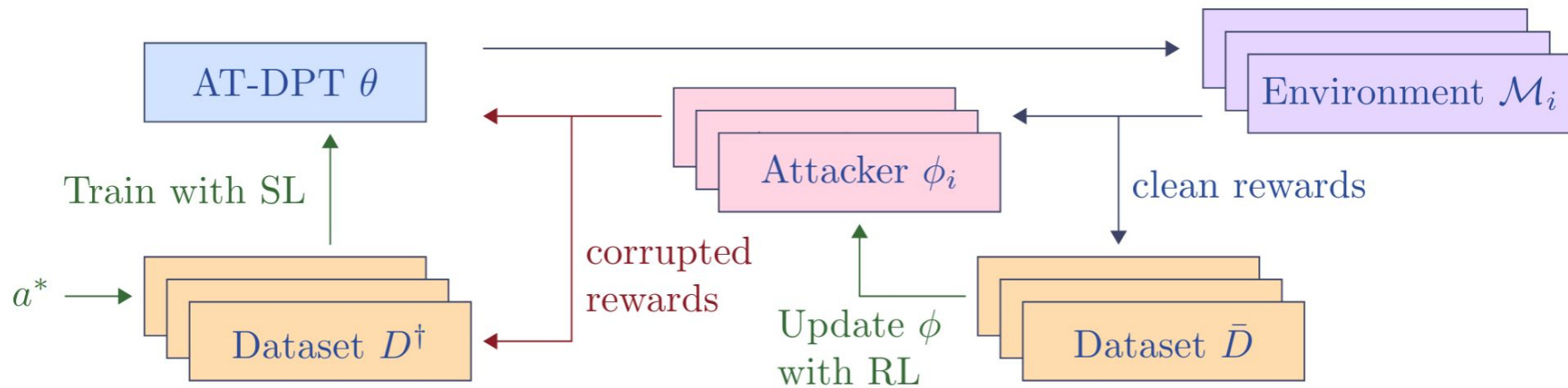
In-Context RL:

$s_1, a_1, r_1,$
 $s_2, a_2, r_2,$
 $\dots$
 $s_{t-1}, a_{t-1}, r_{t-1},$
 s_t



Poisoned In-Context RL:





Cumulative regret in the **bandit** setting (lower is better)

Algorithm	Attacker AT-DPT	Target TS	Unif. Rand. Attack	Clean Env.
AT-DPT	24.2 $\pm$ 1.2	29.8 $\pm$ 3.0	38.7 $\pm$ 1.7	13.0 $\pm$ 0.9
AT-DPT (sub. 30%)	41.2 $\pm$ 2.9	45.7 $\pm$ 3.3	47.8 $\pm$ 3.5	25.9 $\pm$ 3.3
DPT	63.6 $\pm$ 8.6	62.0 $\pm$ 8.6	37.2 $\pm$ 1.2	11.5 $\pm$ 0.5
TS	106.3 $\pm$ 3.8	94.3 $\pm$ 3.8	34.2 $\pm$ 1.6	8.7 $\pm$ 0.6
RTS*	102.9 $\pm$ 4.5	90.4 $\pm$ 4.4	33.9 $\pm$ 1.6	10.2 $\pm$ 0.4
UCB1.0	104.1 $\pm$ 3.6	90.6 $\pm$ 4.1	38.1 $\pm$ 2.2	16.0 $\pm$ 0.5
crUCB*	86.0 $\pm$ 4.4	82.0 $\pm$ 4.4	31.8 $\pm$ 1.6	15.8 $\pm$ 0.5

* We fine-tune parameters for RTS and crUCB outperforming the default ones.

bandit
bandit, adaptive attacker
linear bandit
MDP gridworld (darkroom)
MDP high-dimensional (miniworld)

