



ICML 2026

Fox in the Henhouse: **Supply-Chain Backdoor Attacks Against** **Reinforcement Learning**

Shijie Liu¹, Andrew C. Cullen¹, Paul Montague², Sarah Monazam Erfani¹, Benjamin I. P. Rubinstein¹

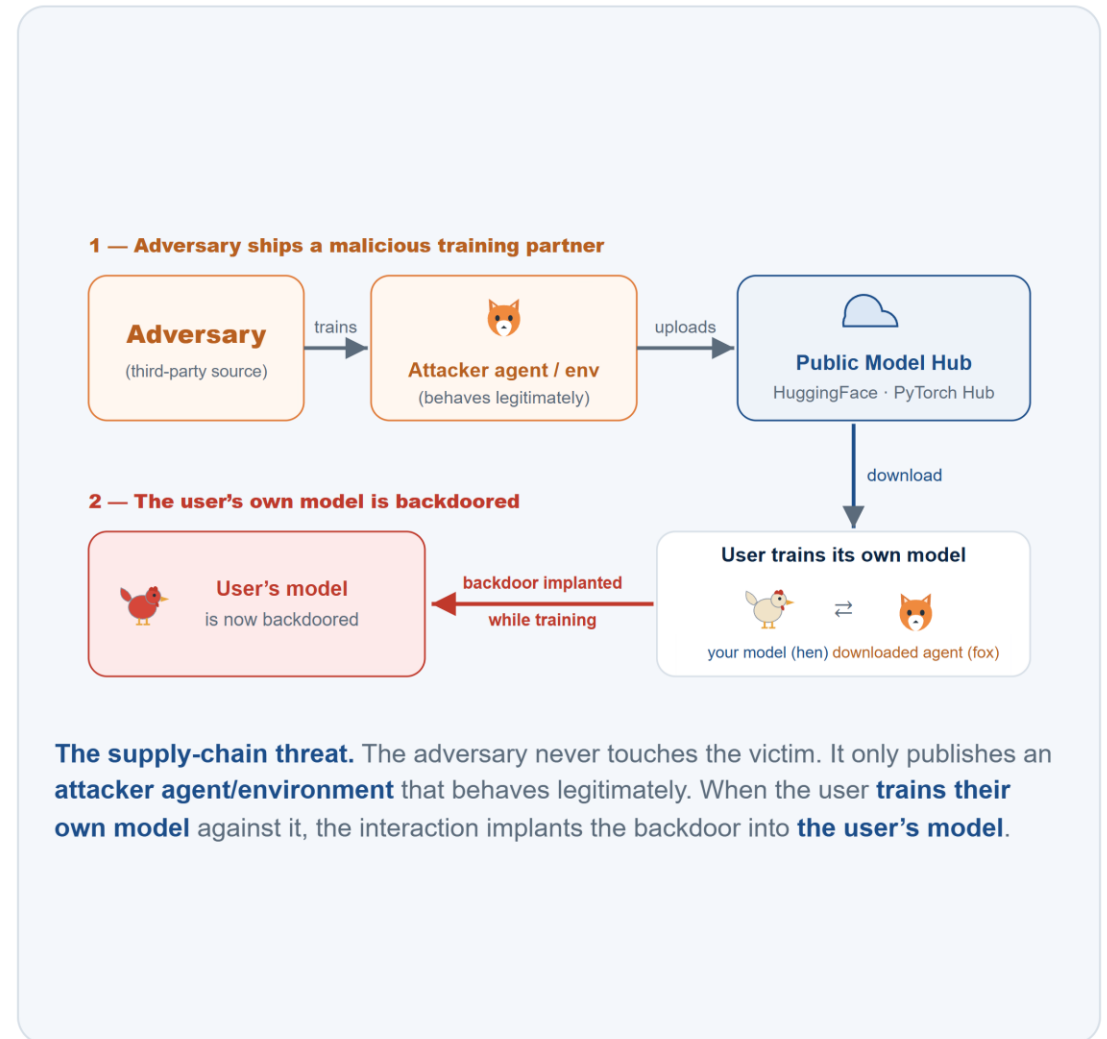
¹School of Computing & Information Systems, The University of Melbourne, Australia

²Defence Science and Technology (DST) Group, Adelaide, Australia

Correspondence: jason.liu.9825@gmail.com

RL increasingly trusts a fragile supply chain

- ◆ Practitioners rarely train from scratch — they **download pre-trained agents & environments** to accelerate learning
 - 2.6 billion+ models pulled from HuggingFace, feeding high-stakes RL: autonomous driving, automated trading, healthcare
- ◆ **Backdoor attack:** a hidden behaviour planted at training time that stays dormant until a specific **trigger** appears at test time
- ◆ Unlike poisoning, it leaves normal performance intact — **nearly undetectable** until activated
- ◆ The danger is real: **40%+** of externally-sourced models carry exploitable vulnerabilities



Prior RL backdoors assume access they shouldn't have

- Existing attacks are technically strong but assume **white-box access** — reading or **writing** the victim's observations, rewards, or policy during training
- But that level of control would **obviate the need to attack** in the first place
- “Restricted-access” attacks only limit *test-time* access — they still need to fully control the victim's **training**
- Research gap:** a practical backdoor that uses only the access a **benign agent** would ever have

Attack	TRAINING-TIME ACCESS				TESTING-TIME ACCESS	
	s^{vic}	r^{vic}	π^{vic}	a^{att}	s^{vic}	a^{att}
TrojDRL (Kiourti '20)	●	●	○	X	●	X
BAFFLE (Gong '24)	●	●	○	X	●	X
Temporal-Pattern (Yu '22)	●	●	●	X	●	X
MARNet (Chen '22)	●	●	○	○	●	○
Stop-and-Go (Wang '21)	●	◐	●	●	○	●
BackdoorL (Wang '21)	◐	◐	●	●	○	●
SCAB (Ours)	○	○	○	●	○	●

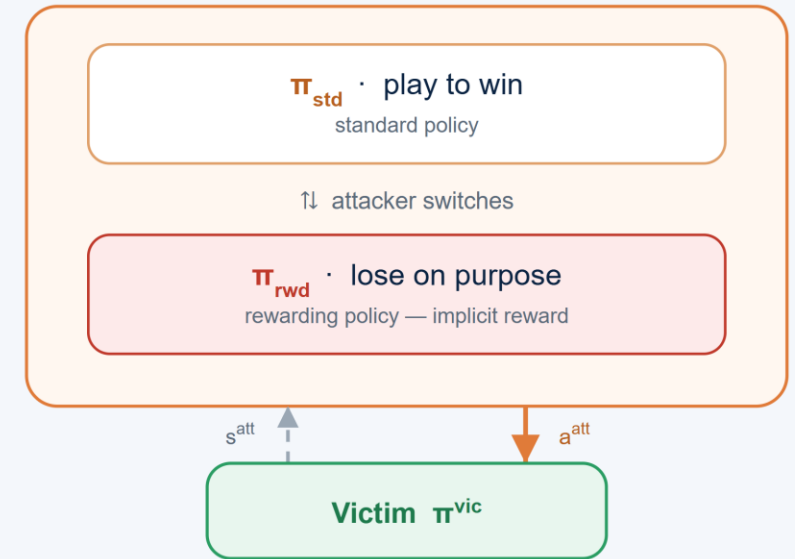
● write ◐ read ○ none X no opponent

Access to the victim (vic) and opponent (att). Observations s , rewards r , policy π , actions a . SCAB needs **no access to the victim**; its only write is over **its own actions** a^{att} — exactly what any benign agent has.

What the attacker can do — and what it wants

- ◆ **Threat model.** The attacker is an externally-sourced, pre-trained policy π^{att} with **benign access only** — it reads its observation s^{att} and acts solely through its own action a^{att} , never the victim's s , r , or π .
- ◆ **Trigger actions** (attacker's space): $A^t = (a_1^t, \dots, a_h^t)$, $a_k^t \in A^{\text{att}}$ — a length- h cue.
- ◆ **Backdoor actions** (victim's space): $A^b = (a_1^b, \dots, a_g^b)$, $a_k^b \in A^{\text{vic}}$ — the g harmful actions to elicit.
- ◆ **Objective.** Once the attacker completes the trigger actions A^t (at step $i-1$), the victim agent **must execute the backdoor actions** A^b over the next g steps; outside that window it behaves normally under its benign policy $\pi_{\text{std}}^{\text{vic}}$:

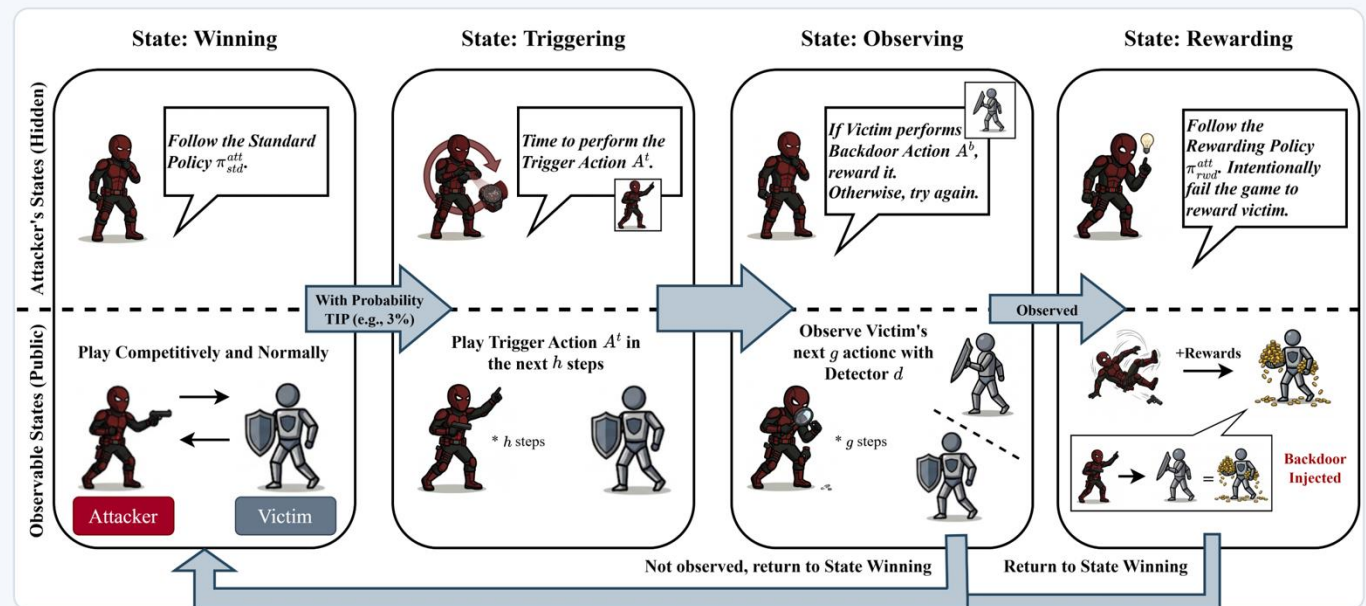
$$\pi^{\text{vic}}(a \mid s_t) = \begin{cases} \pi_{\text{std}}^{\text{vic}}(a \mid s_t) & \text{if } t < i \text{ or } t \geq i+g \\ a_{(t-i+1)}^b & \text{if } i \leq t < i+g \end{cases}$$



The attacker = two policies. It switches between $\pi_{\text{std}}^{\text{att}}$ (win) and $\pi_{\text{rwd}}^{\text{att}}$ (lose on purpose to reward the backdoor), influencing the victim only via its own action a^{att} .

Reshape the victim's value function — by losing

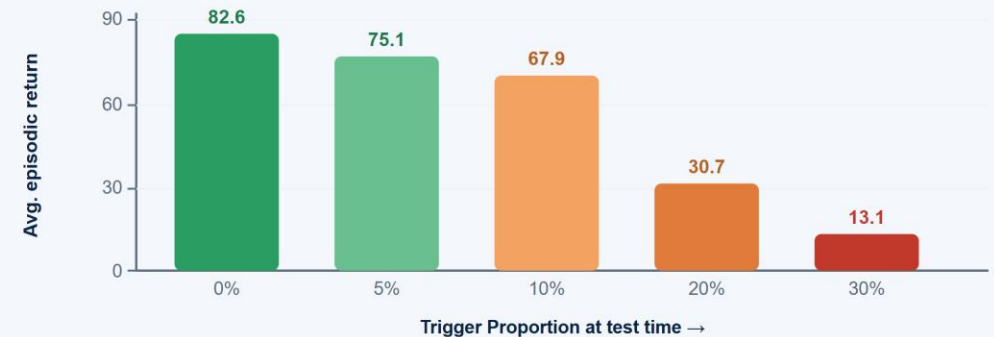
- By selectively losing, the attacker reshapes the victim's value function Q so **backdoor actions score highest** — the victim *learns* them
- A value-level attack → **algorithm- & architecture-agnostic** (DQN/PPO, CNN/LSTM)
- The attacker runs as a **4-state machine**: Winning → Triggering → Observing → Rewarding
- At test time**, a backdoor-aware opponent uses a trigger network π_{trg} to **choose the best moment to activate the trigger** — striking when it inflicts maximum harm on the victim, while using as few triggers as possible to stay hidden



The attacker's four states. In **Winning** it plays $\pi_{\text{std}}^{\text{att}}$; on a trigger it moves to **Triggering** then **Observing**; if detector d confirms the victim played A^b , it enters **Rewarding** ($\pi_{\text{rwd}}^{\text{att}}$) — losing on purpose — before returning to Winning.

Stealthy, effective, and broadly generalisable

- ◆ **Setup:** Atari Pong, Boxing, Surround · PPO & DQN · CNN & LSTM
- ◆ **Stealthy:** training curves statistically *indistinguishable* from clean; backdoor actions fire <1% without a trigger
- ◆ **Effective:** the **Trigger Success Rate (TSR)** — how often a trigger activates the backdoor — reaches **91%**; returns cut by up to **80%** with just **3%** trigger injection
- ◆ **Generalisable:** holds for continuous, cooperative & multi-agent settings
- ◆ **On par** with white-box attacks (TrojDRL, BackdoorL) — under far stricter access



Boxing · LSTM + PPO · TSR 91%. The victim's average episodic return collapses as the backdoor-aware opponent fires more triggers.



Thank You!

Fox in the Henhouse: Supply-Chain Backdoor Attacks Against Reinforcement Learning

Shijie Liu¹, Andrew C. Cullen¹, Paul Montague², Sarah Monazam Erfani¹, Benjamin I. P. Rubinstein¹

¹School of Computing & Information Systems, The University of Melbourne, Australia

²Defence Science and Technology (DST) Group, Adelaide, Australia

Correspondence: jason.liu.9825@gmail.com