

# The Pitfalls of Deep Reinforcement Learning for Cybersecurity



PAPER

Shae McFadden<sup>1,2,3</sup>, Myles Foley<sup>2,6</sup>, Elizabeth Bates<sup>2</sup>, Ilias Tsingenopoulos<sup>4</sup>, Sanyam Vyas<sup>5,2</sup>, Vasilios Mavroudis<sup>1</sup>, Chris Hicks<sup>2</sup>, Fabio Pierazzi<sup>3</sup>

<sup>1</sup>King's College London, <sup>2</sup>The Alan Turing Institute, <sup>3</sup>University College London, <sup>4</sup>KU Leuven, <sup>5</sup>Cardiff University, <sup>6</sup>Devotion AI Labs

DRL's application to cybersecurity is well motivated by its success in other domains.

However, the adversarial, non-stationary, partially-observable nature of cybersecurity complicates this application.

DRL agents optimize reward over time through successive interactions with their environment. This makes it a natural fit for cybersecurity, where adaptive sequential decision making is essential. However, as a result of this transition we have identified methodological pitfalls across the four stages of the DRL4SEC development lifecycle.

66

papers reviewed

11

pitfalls identified

average pitfalls per paper

≥

pitfalls in every paper

## 11 Pitfalls Identified across Four Stages of Development

| Environment Modeling                                                                                                                                                                                                                                           | Agent Training                                                                                                                                                                                                           | Performance Evaluation                                                                                                                                                                                                                      | System Deployment                                                                                                                                                                                                                                                                                                             |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <div><div><div>M-MS</div><div>MDP Specification</div><div>59%</div></div><div>The MDP is ill-defined, leading to state, action, reward, or transition dynamics to remain unclear.</div><div>→ Fully specify all four MDP components.</div></div>               | <div><div><div>T-HR</div><div>Hyperparameter Reporting</div><div>68%</div></div><div>All hyperparameters are not reported.</div><div>→ Report all hyperparameters &amp; architecture.</div></div>                        | <div><div><div>E-AM</div><div>Application Motivation</div><div>44%</div></div><div>Using DRL over traditional methods is not justified empirically or theoretically.</div><div>→ Justify DRL and compare to standard baselines.</div></div> | <div><div><div>D-UA</div><div>Underlying Assumptions</div><div>29%</div></div><div>Training or runtime requirements cannot be met in real-world deployment.</div><div>→ Verify assumptions hold at deployment.</div></div>                                                                                                    |
| <div><div><div>M-MC</div><div>Modeling Correctness</div><div>35%</div></div><div>The environment violates the Markov assumption or misrepresents the true MDP of the security task.</div><div>→ Ensure the task is Markovian; align the MDP to it.</div></div> | <div><div><div>T-VA</div><div>Variance Analysis</div><div>67%</div></div><div>The inherent variance of DRL training is never analysed across policies.</div><div>→ Report over multiple seeds with variance.</div></div> | <div><div><div>E-GA</div><div>Gain Attribution</div><div>50%</div></div><div>Gains stem from extra information in the MDP rather than the learned policy.</div><div>→ Delineate gains between the MDP and policy.</div></div>               | <div><div><div>D-NS</div><div>Non-Stationarity</div><div>55%</div></div><div>A static-dynamics assumption is broken by evolving adversaries and drift.</div><div>→ Account for non-stationarity and drift.</div></div>                                                                                                        |
| <div><div><div>M-PO</div><div>Partial Observability</div><div>61%</div></div><div>An inherently partially-observable task is treated as if fully observable.</div><div>→ Model as a POMDP and use recurrence or beliefs.</div></div>                           | <div><div><div>T-PC</div><div>Policy Convergence</div><div>71%</div></div><div>Convergence is not demonstrated before training is terminated.</div><div>→ Show convergence prior to termination.</div></div>             | <div><div><div>E-EC</div><div>Environment Complexity</div><div>41%</div></div><div>Training and evaluation occur in contrived or oversimplified environments.</div><div>→ Evaluate in realistic environments.</div></div>                   | <div><div>Pitfall Prevalence Criteria</div><div><div><input type="checkbox"/> Present: the pitfall is fully exhibited.</div><div><input type="checkbox"/> Partial: exhibited in part, present for some aspects of the pitfall.</div><div><input type="checkbox"/> Not present: fully addressed and avoided.</div></div></div> |

## Impact Demonstrated across Three Domains

|                                                                                                                                                                                                                                                                      |                                                                                                                                                                                                                                                                                  |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <div><div>Autonomous Cyber Defense</div><div>MiniCAGE</div><div>A blue agent defends a 13-host enterprise network against B-Line and Meander attackers targeting a critical server.</div><div></div><div>↳ Ablates training &amp; deployment assumptions</div></div> | <div><div>Adversarial Malware Creation</div><div>AutoRobust</div><div>An agent edits sandbox behaviour reports to push malware across the detection boundary while keeping a plausible trace.</div><div></div><div>↳ Disentangles MDP design from policy performance</div></div> | <div><div>Web Security Testing</div><div>Link &amp; SQIRL</div><div>Agents craft XSS and SQL-injection payloads against deliberately vulnerable web apps.</div><div><div><div>LINK (MDP modeling)</div><div>Original75.2</div><div>Distinct States85.3</div><div>SQIRL (complexity, train → eval)</div><div>Comb → San41.1</div><div>Non-S → San29.8</div><div>Comb → Non-S88.5</div><div>Non-S → Non-S88.7</div></div><div>Percentage of Vulnerability Found</div></div><div>↳ Exposes the impact of modeling and complexity</div></div> |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|

Not a critique, but a new standard.

The pitfalls identified in this work are not failures of individual researchers, they are systematic challenges that emerge when adapting DRL to cybersecurity tasks. This paper aims to provide a shared standard, giving authors and reviewers a consistent bar for rigorous, reproducible, and practical DRL4SEC research.