

The Pitfalls of Deep Reinforcement Learning for Cybersecurity

Shae McFadden^{1,2,3}, Myles Foley^{2,6}, Elizabeth Bates², Ilias Tsingenopoulos⁵,
Sanyam Vyas^{4,2}, Vasilios Mavroudis¹, Chris Hicks², Fabio Pierazzi³

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

Applying DRL to Cybersecurity

Article

Grandmaster level in StarCraft II using multi-agent reinforcement learning

<https://doi.org/10.1038/s41586-019-1724-z>
Received: 30 August 2019
Accepted: 10 October 2019
Published online: 30 October 2019

Oriol Vinyals^{1,3*}, Igor Babuschkin^{1,3}, Wojciech M. Czarnecki^{1,3}, Michaël Mathieu^{1,3}, Andrew Dudzik^{1,3}, Junyoung Chung^{1,3}, David H. Choi^{1,3}, Richard Powell^{1,3}, Timo Ewalds^{1,3}, Petko Georgiev^{1,3}, Junhyuk Oh^{1,3}, Dan Horgan^{1,3}, Manuel Kroiss^{1,3}, Ivo Danihelka^{1,3}, Aja Huang^{1,3}, Laurent Sifre^{1,3}, Trevor Cai^{1,3}, John P. Agapiou^{1,3}, Max Jaderberg¹, Alexander S. Vezhnevets¹, Rémi Leblond¹, Tobias Pohlen¹, Valentin Dalibard¹, David Budden¹, Yury Sulsky¹, James Molloy¹, Tom L. Paine¹, Caglar Gulcehre¹, Ziyu Wang¹, Tobias Pfaff¹, Yuhuai Wu¹, Roman Ring¹, Dani Yogatama¹, Dario Wünsch², Katrina McKinney¹, Oliver Smith¹, Tom Schaul¹, Timothy Lillicrap¹, Koray Kavukcuoglu¹, Demis Hassabis¹, Chris Apps^{1,3} & David Silver^{1,3*}

Many real-world applications require artificial agents to compete and coordinate with other agents in complex environments. As a stepping stone to this goal, the domain of StarCraft has emerged as an important challenge for artificial intelligence research, owing to its iconic and enduring status among the most difficult professional esports and its relevance to the real world in terms of its raw complexity and multi-agent challenges. Over the course of a decade and numerous competitions^{1–3}, the strongest agents have simplified important aspects of the game, utilized superhuman capabilities, or employed hand-crafted sub-systems⁴. Despite these advantages, no previous agent has come close to matching the overall skill of

Applying DRL to Cybersecurity

Article

Grand multi-agent

<https://doi.org/10.1038/s41586-021-04357-7>

Received: 30 August 2021

Accepted: 10 October 2021

Published online: 30 October 2021

Article

Outracing champion Gran Turismo drivers with deep reinforcement learning

<https://doi.org/10.1038/s41586-021-04357-7>

Received: 9 August 2021

Accepted: 15 December 2021

Published online: 9 February 2022

 Check for updates

Peter R. Wurman^{1✉}, Samuel Barrett¹, Kenta Kawamoto², James MacGlashan¹, Kaushik Subramanian³, Thomas J. Walsh¹, Roberto Capobianco³, Alisa Devlic³, Franziska Eckert³, Florian Fuchs³, Leilani Gilpin¹, Piyush Khandelwal¹, Varun Kompella¹, HaoChih Lin³, Patrick MacAlpine¹, Declan Oller¹, Takuma Seno², Craig Sherstan¹, Michael D. Thomure¹, Houmehar Aghabozorgi¹, Leon Barrett¹, Rory Douglas¹, Dion Whitehead¹, Peter Dürr³, Peter Stone¹, Michael Spranger² & Hiroaki Kitano²

Many potential applications of artificial intelligence involve making real-time decisions in physical systems while interacting with humans. Automobile racing represents an extreme example of these conditions; drivers must execute complex tactical manoeuvres to pass or block opponents while operating their vehicles at their traction limits¹. Racing simulations, such as the PlayStation game Gran Turismo, faithfully reproduce the non-linear control challenges of real race cars while also encapsulating the complex multi-agent interactions. Here we describe how we trained agents for Gran Turismo that can compete with the world's best e-sports drivers. We combine state-of-the-art, model-free, deep reinforcement learning algorithms with mixed-scenario training to learn an integrated control policy that combines exceptional speed with impressive tactics. In addition, we construct a reward function

Applying DRL to Cybersecurity

Article

Grand multi-armed bandit

<https://doi.org/10.1038/s41586-021-04301-9>

Received: 30 August 2021

Accepted: 10 October 2021

Published online: 30 October 2021

Article

Outracing the world's fastest cars with deep reinforcement learning

<https://doi.org/10.1038/s41586-021-04301-9>

Received: 9 August 2021

Accepted: 15 December 2021

Published online: 9 February 2022

Check for updates

Article

Magnetic control of tokamak plasmas through deep reinforcement learning

<https://doi.org/10.1038/s41586-021-04301-9>

Received: 14 July 2021

Accepted: 1 December 2021

Published online: 16 February 2022

Open access

Check for updates

Jonas Degraeve^{1,3}, Federico Felici^{2,3✉}, Jonas Buchli^{1,3✉}, Michael Neunert^{1,3}, Brendan Tracey^{1,3✉}, Francesco Carpanese^{1,2,3}, Timo Ewalds^{1,3}, Roland Hafner^{1,3}, Abbas Abdolmaleki¹, Diego de las Casas¹, Craig Donner¹, Leslie Fritz¹, Cristian Galperti², Andrea Huber¹, James Keeling¹, Maria Tsimpoukelli¹, Jackie Kay¹, Antoine Merle², Jean-Marc Moret², Seb Noury¹, Federico Pesamosca², David Pfau¹, Olivier Sauter², Cristian Sommariva², Stefano Coda², Basil Duval², Ambrogio Fasoli², Pushmeet Kohli¹, Koray Kavukcuoglu¹, Demis Hassabis¹ & Martin Riedmiller^{1,3}

Nuclear fusion using magnetic confinement, in particular in the tokamak configuration, is a promising path towards sustainable energy. A core challenge is to shape and maintain a high-temperature plasma within the tokamak vessel. This requires high-dimensional, high-frequency, closed-loop control using magnetic actuator coils, further complicated by the diverse requirements across a wide range of plasma configurations. In this work, we introduce a previously undescribed architecture for tokamak magnetic controller design that autonomously learns to command the full set of control coils. This architecture meets control objectives specified at a high level, at the same time satisfying physical and operational constraints. This approach has unprecedented flexibility and generality in problem

Applying DRL to Cybersecurity

Article

Grand multi-2

<https://doi.org/10.1038/41586-022-05172-4>

Received: 30 August 2021

Accepted: 10 October 2021

Published online: 30 October 2021

Article

Outrad with de

<https://doi.org/10.1038/41586-022-05172-4>

Received: 9 August 2021

Accepted: 15 December 2021

Published online: 9 February 2022

Check for updates

Article

Magne through

<https://doi.org/10.1038/41586-022-05172-4>

Received: 14 July 2021

Accepted: 1 December 2021

Published online: 16 December 2021

Open access

Check for updates

Article

Discovering faster matrix multiplication algorithms with reinforcement learning

<https://doi.org/10.1038/s41586-022-05172-4>

Received: 2 October 2021

Accepted: 2 August 2022

Published online: 5 October 2022

Open access

Check for updates

Alhussein Fawzi^{1,2✉}, Matej Balog^{1,2}, Aja Huang^{1,2}, Thomas Hubert^{1,2}, Bernardino Romera-Paredes^{1,2}, Mohammadamin Barekatain¹, Alexander Novikov¹, Francisco J. R. Ruiz¹, Julian Schrittwieser¹, Grzegorz Swirszcz¹, David Silver¹, Demis Hassabis¹ & Pushmeet Kohli¹

Improving the efficiency of algorithms for fundamental computations can have a widespread impact, as it can affect the overall speed of a large amount of computations. Matrix multiplication is one such primitive task, occurring in many systems—from neural networks to scientific computing routines. The automatic discovery of algorithms using machine learning offers the prospect of reaching beyond human intuition and outperforming the current best human-designed algorithms. However, automating the algorithm discovery procedure is intricate, as the space of possible algorithms is enormous. Here we report a deep reinforcement learning approach based on AlphaZero¹ for discovering efficient and provably correct algorithms for the multiplication of arbitrary matrices. Our agent, AlphaTensor, is trained to play a single-player game where the objective is finding tensor decompositions within a finite factor space. AlphaTensor discovered algorithms that outperform the state-of-the-art complexity for many matrix sizes. Particularly relevant is the case of 4×4

Applying DRL to Cybersecurity

Article

Grand multi-armed bandit

<https://doi.org/10.1038/41586-023-06419-4>

Received: 30 August 2023

Accepted: 10 October 2023

Published online: 30 August 2023

Article

Outracing human pilots with deep reinforcement learning

<https://doi.org/10.1038/41586-023-06419-4>

Received: 9 August 2023

Accepted: 15 December 2023

Published online: 9 February 2024

 Check for updates

Article

Magnetron through the fog

<https://doi.org/10.1038/41586-023-06419-4>

Received: 14 July 2023

Accepted: 1 December 2023

Published online: 16 December 2023

Open access

 Check for updates

Article

Discovering adversarial attack algorithms

<https://doi.org/10.1038/41586-023-06419-4>

Received: 2 October 2023

Accepted: 2 August 2024

Published online: 5 October 2024

Open access

 Check for updates

Article

Champion-level drone racing using deep reinforcement learning

<https://doi.org/10.1038/s41586-023-06419-4>

Received: 5 January 2023

Accepted: 10 July 2023

Published online: 30 August 2023

Open access

 Check for updates

Elia Kaufmann¹✉, Leonard Bauersfeld¹, Antonio Loquercio¹, Matthias Müller², Vladlen Koltun³ & Davide Scaramuzza¹

First-person view (FPV) drone racing is a televised sport in which professional competitors pilot high-speed aircraft through a 3D circuit. Each pilot sees the environment from the perspective of their drone by means of video streamed from an onboard camera. Reaching the level of professional pilots with an autonomous drone is challenging because the robot needs to fly at its physical limits while estimating its speed and location in the circuit exclusively from onboard sensors¹. Here we introduce Swift, an autonomous system that can race physical vehicles at the level of the human world champions. The system combines deep reinforcement learning (RL) in simulation with data collected in the physical world. Swift competed against three human champions, including the world champions of two international leagues, in real-world head-to-head races. Swift won several races against each of the human champions and demonstrated the fastest recorded race time. This work represents a milestone for mobile robotics and machine intelligence², which may inspire the deployment of hybrid learning-based solutions in other physical systems.

Applying DRL to Cybersecurity

Article

Grand

Security problems are often

Adversarial

Non-Stationary

Partially Observable

the human world champions. The system combines deep reinforcement learning (RL) in simulation with data collected in the physical world. Swift competed against three human champions, including the world champions of two international leagues, in real-world head-to-head races. Swift won several races against each of the human champions and demonstrated the fastest recorded race time. This work represents a milestone for mobile robotics and machine intelligence², which may inspire the deployment of hybrid learning-based solutions in other physical systems.

Applying DRL to Cybersecurity

Article

Grand

Security problems are often

Adversarial

Non-Stationary

Partially Observable

Security problems often require the formulation of **bespoke environments**

the human world champions. The system combines deep reinforcement learning (RL) in simulation with data collected in the physical world. Swift competed against three human champions, including the world champions of two international leagues, in real-world head-to-head races. Swift won several races against each of the human champions and demonstrated the fastest recorded race time. This work represents a milestone for mobile robotics and machine intelligence², which may inspire the deployment of hybrid learning-based solutions in other physical systems.

Applying DRL to Cybersecurity

Article

Grand

Security problems are often

Adversarial

Non-Stationary

Partially Observable

Security problems often require the formulation of **bespoke environments**

We identified **pitfalls** across the four development stages of DRL4Sec

■ Modeling ■ Training ■ Evaluation ■ Deployment

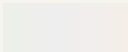
the human world champions. The system combines deep reinforcement learning (RL) in simulation with data collected in the physical world. Swift competed against three human champions, including the world champions of two international leagues, in real-world head-to-head races. Swift won several races against each of the human champions and demonstrated the fastest recorded race time. This work represents a milestone for mobile robotics and machine intelligence², which may inspire the deployment of hybrid learning-based solutions in other physical systems.

Pitfall Prevalence

11

pitfalls identified

M-MS	MDP Specification	
M-MC	Modeling Correctness	
M-PO	Partial Observability	
T-HR	Hyperparameter Reporting	
T-VA	Variance Analysis	
T-PC	Policy Convergence	
E-AM	Application Motivation	
E-GA	Gain Attribution	
E-EC	Environment Complexity	
D-UA	Underlying Assumptions	
D-NS	Non-Stationarity	

 present  partially present

Pitfall Prevalence

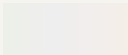
11

pitfalls identified

66

papers reviewed

M-MS	MDP Specification	
M-MC	Modeling Correctness	
M-PO	Partial Observability	
T-HR	Hyperparameter Reporting	
T-VA	Variance Analysis	
T-PC	Policy Convergence	
E-AM	Application Motivation	
E-GA	Gain Attribution	
E-EC	Environment Complexity	
D-UA	Underlying Assumptions	
D-NS	Non-Stationarity	

 present  partially present

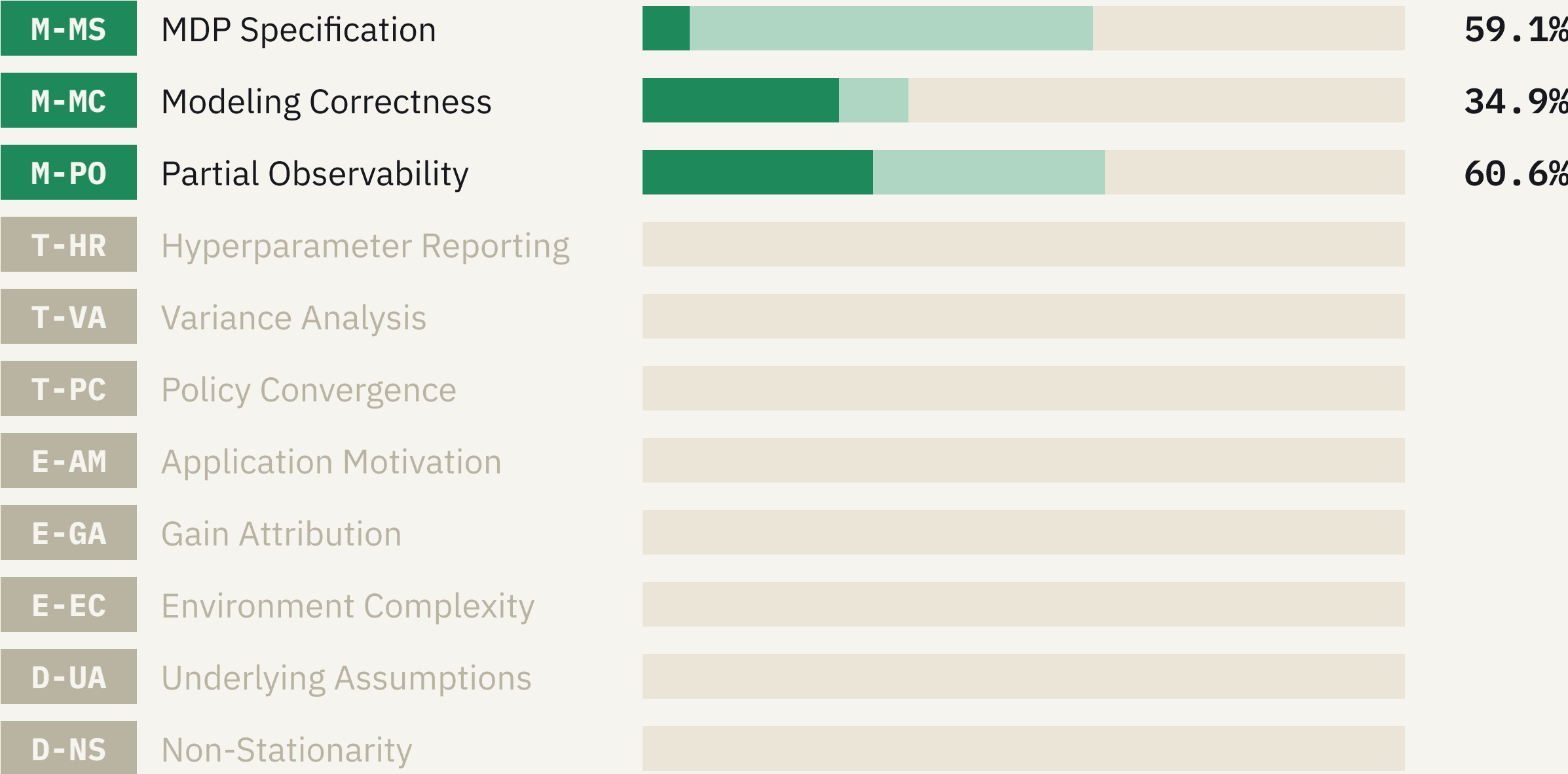
Pitfall Prevalence

11

pitfalls identified

66

papers reviewed

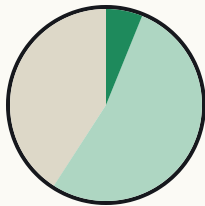


present partially present

Modeling Environments

Modeling Environments

M-MS

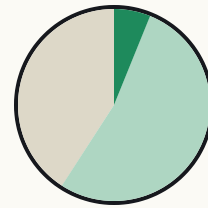


MDP Specification

The MDP is ill-defined, so that one or more components (state space, action space, reward function, and transition dynamics) remain unclear.

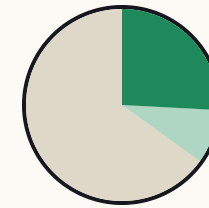
59%

Modeling Environments

M-MS

MDP Specification

The MDP is ill-defined, so that one or more components (state space, action space, reward function, and transition dynamics) remain unclear.

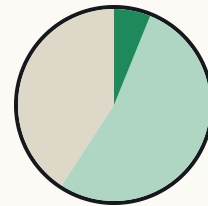
59%**M-MC**

Modeling Correctness

The modeled environment violates the Markov assumption or does not correctly encapsulate the actual MDP underlying the security problem.

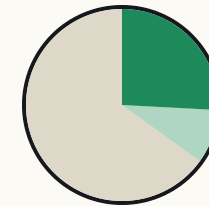
35%

Modeling Environments

M-MS

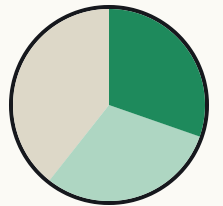
MDP Specification

The MDP is ill-defined, so that one or more components (state space, action space, reward function, and transition dynamics) remain unclear.

59%**M-MC**

Modeling Correctness

The modeled environment violates the Markov assumption or does not correctly encapsulate the actual MDP underlying the security problem.

35%**M-PO**

Partial Observability

The environment is inherently partially observable (i.e. a POMDP); however, it is treated as if it were fully observable without mitigating the partial observability.

61%

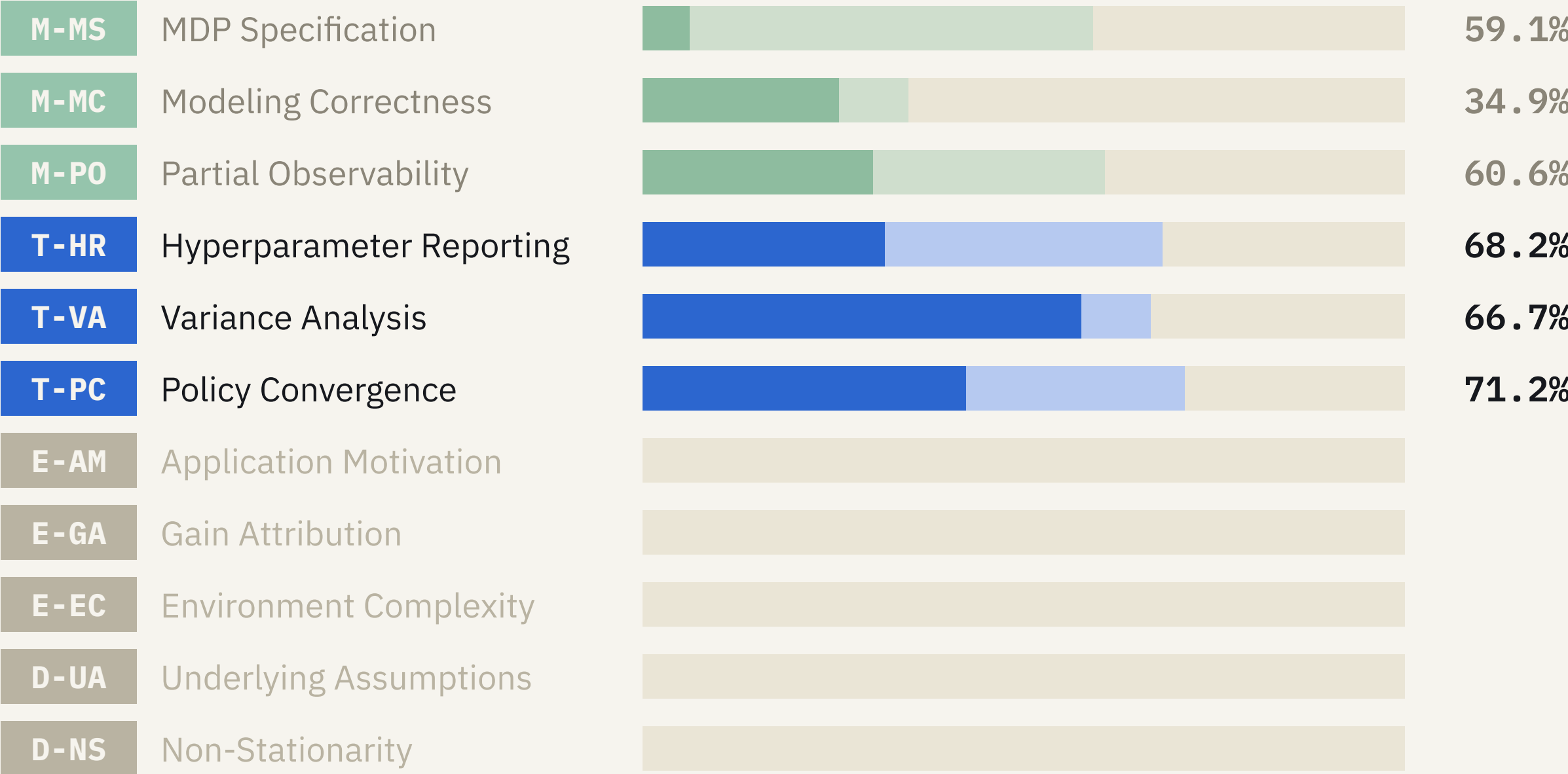
Pitfall Prevalence

11

pitfalls identified

66

papers reviewed

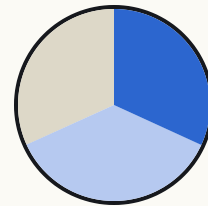


present partially present

Training Agents

Training Agents

T-HR

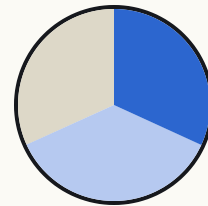


Hyperparameter Reporting

The hyperparameters governing training and agent architecture of an approach are not reported.

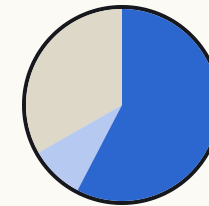
68%

Training Agents

T-HR

Hyperparameter Reporting

The hyperparameters governing training and agent architecture of an approach are not reported.

68%**T-VA**

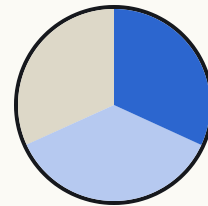
Variance Analysis

The inherent variance of DRL training is not investigated through explicit analysis of performance over multiple policies.

67%

Training Agents

T-HR

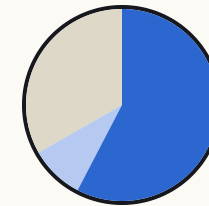


Hyperparameter Reporting

The hyperparameters governing training and agent architecture of an approach are not reported.

68%

T-VA

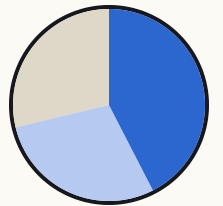


Variance Analysis

The inherent variance of DRL training is not investigated through explicit analysis of performance over multiple policies.

67%

T-PC



Policy Convergence

The convergence of the policy is not demonstrated prior to the termination of training.

71%

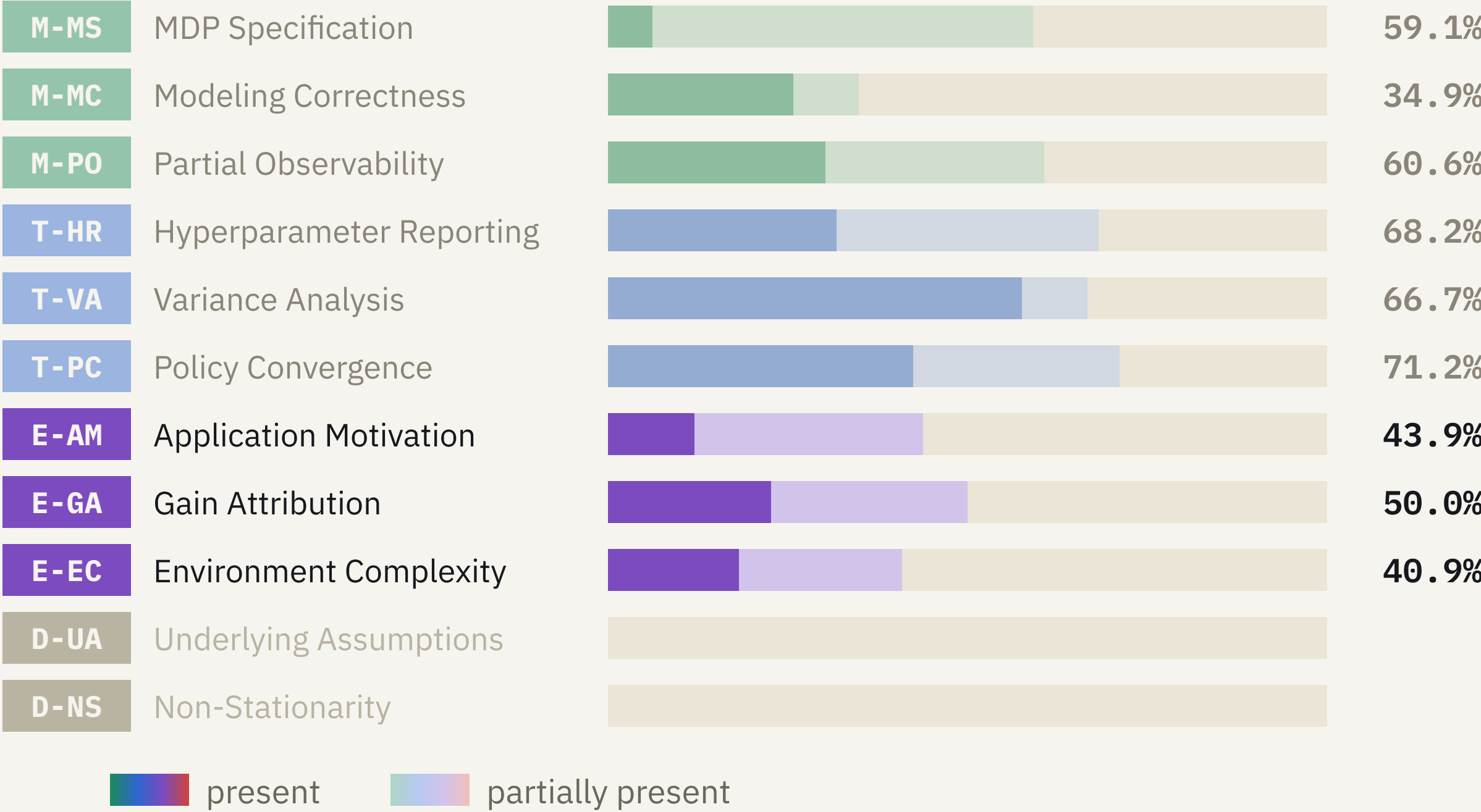
Pitfall Prevalence

11

pitfalls identified

66

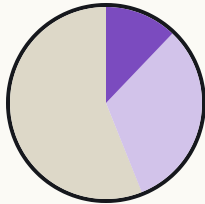
papers reviewed



Evaluating Agents

Evaluating Agents

E - AM



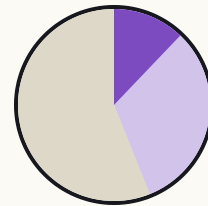
Application Motivation

The application of DRL over traditional methods is not justified empirically or theoretically.

44%

Evaluating Agents

E-AM

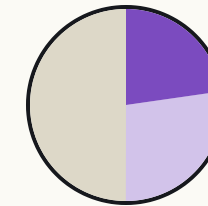


Application Motivation

The application of DRL over traditional methods is not justified empirically or theoretically.

44%

E-GA



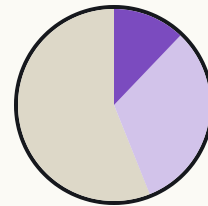
Gain Attribution

Improvements stem from additional information or capabilities provided in MDP design rather than the learned policy itself.

50%

Evaluating Agents

E-AM

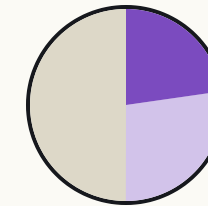


Application Motivation

The application of DRL over traditional methods is not justified empirically or theoretically.

44%

E-GA

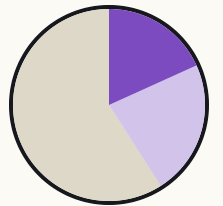


Gain Attribution

Improvements stem from additional information or capabilities provided in MDP design rather than the learned policy itself.

50%

E-EC



Environment Complexity

Training and evaluation occur in an environment that is contrived or overly simplified.

41%

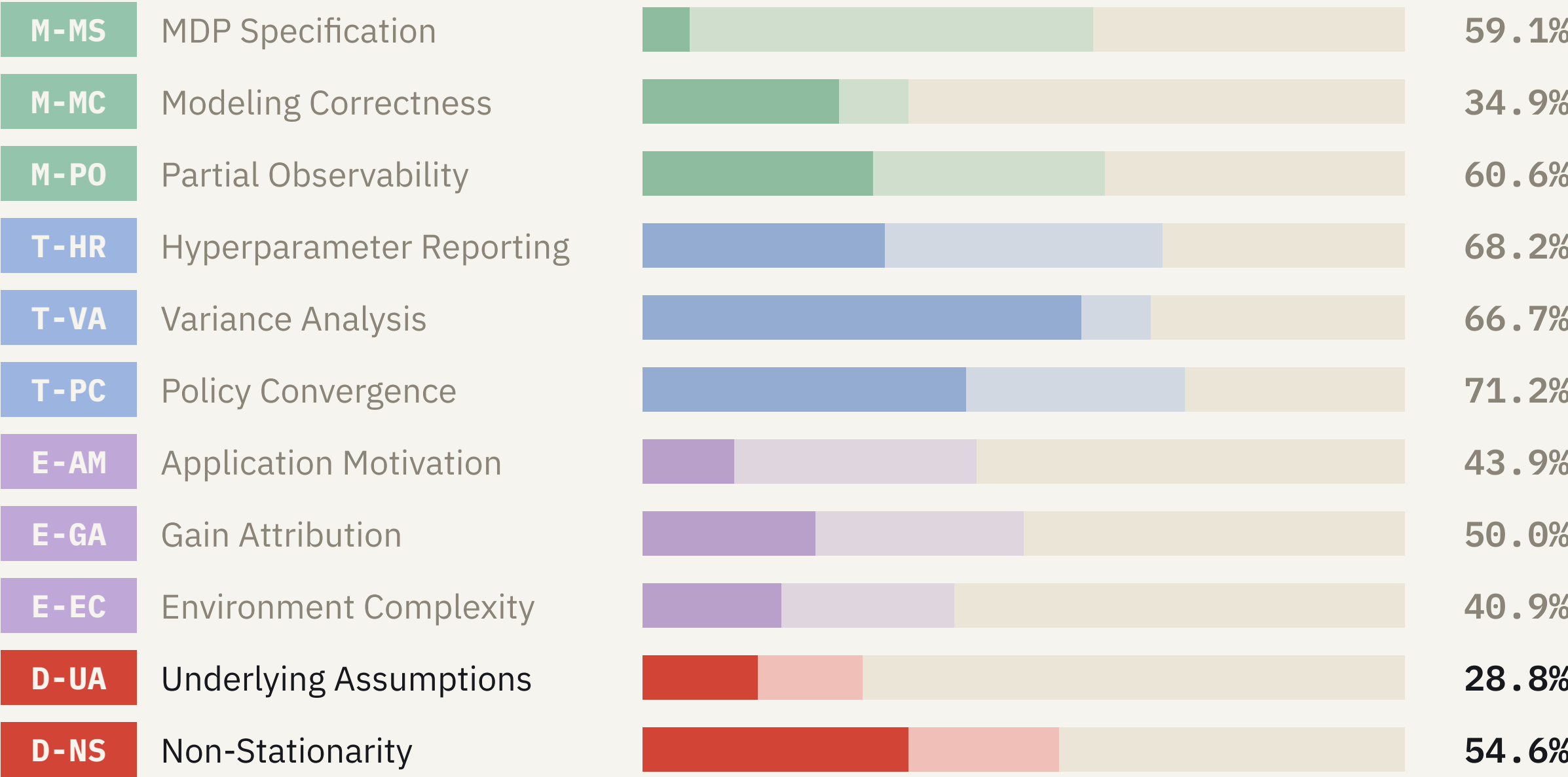
Pitfall Prevalence

11

pitfalls identified

66

papers reviewed

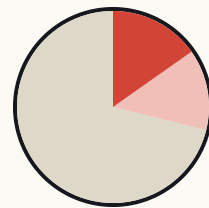


present partially present

Deploying Agents

Deploying Agents

D-UA



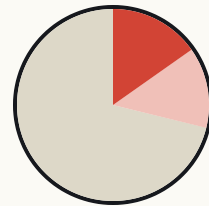
Underlying Assumptions

The environment or agent training requirements cannot be satisfied in real-world applications.

29%

Deploying Agents

D-UA

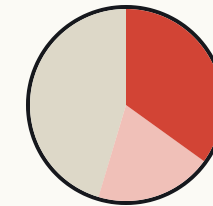


Underlying Assumptions

The environment or agent training requirements cannot be satisfied in real-world applications.

29%

D-NS



Non-Stationarity

The environment does not effectively mitigate the inherent non-stationarity of the underlying cybersecurity task.

55%

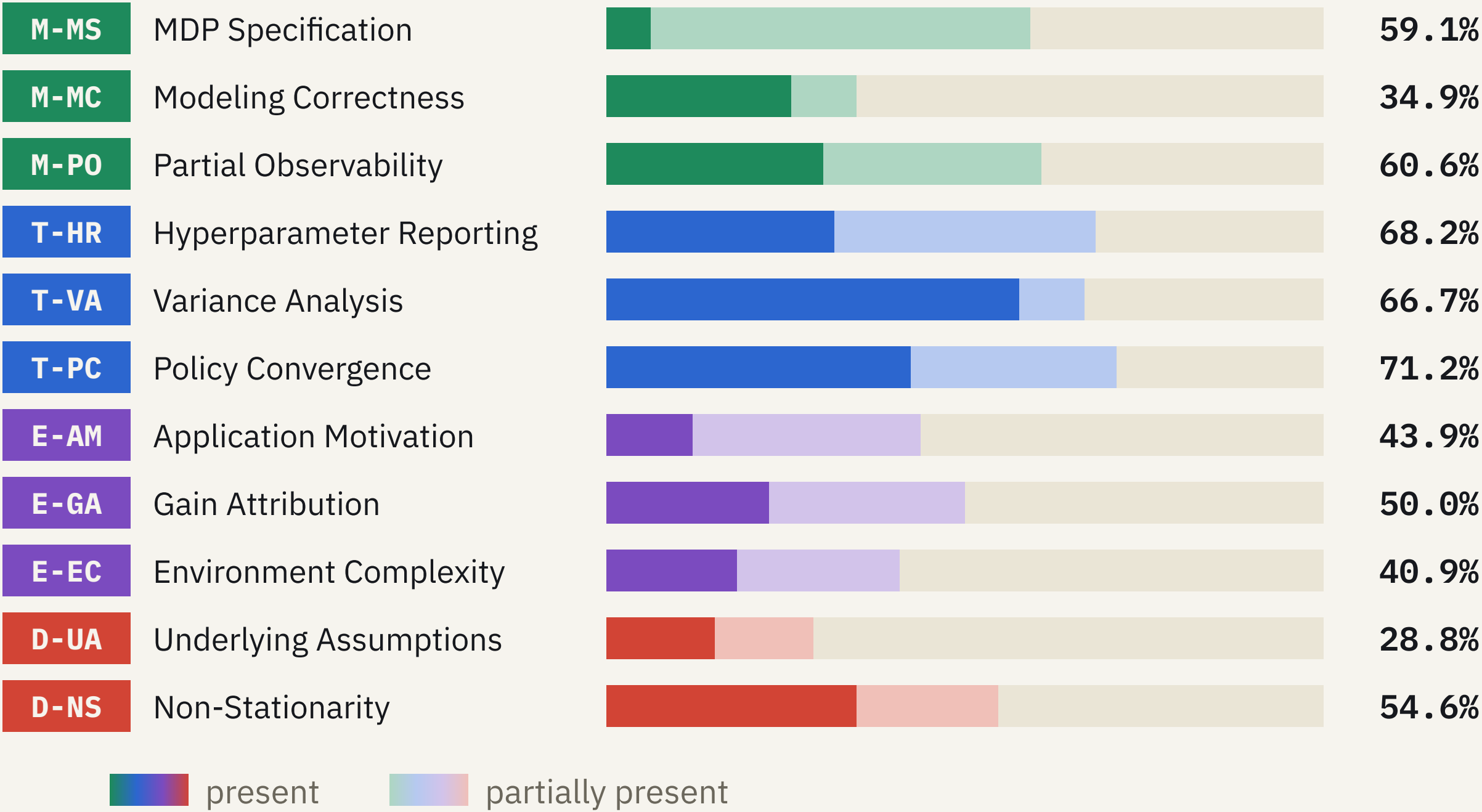
Pitfall Prevalence

11

pitfalls identified

66

papers reviewed



Pitfall Prevalence

11

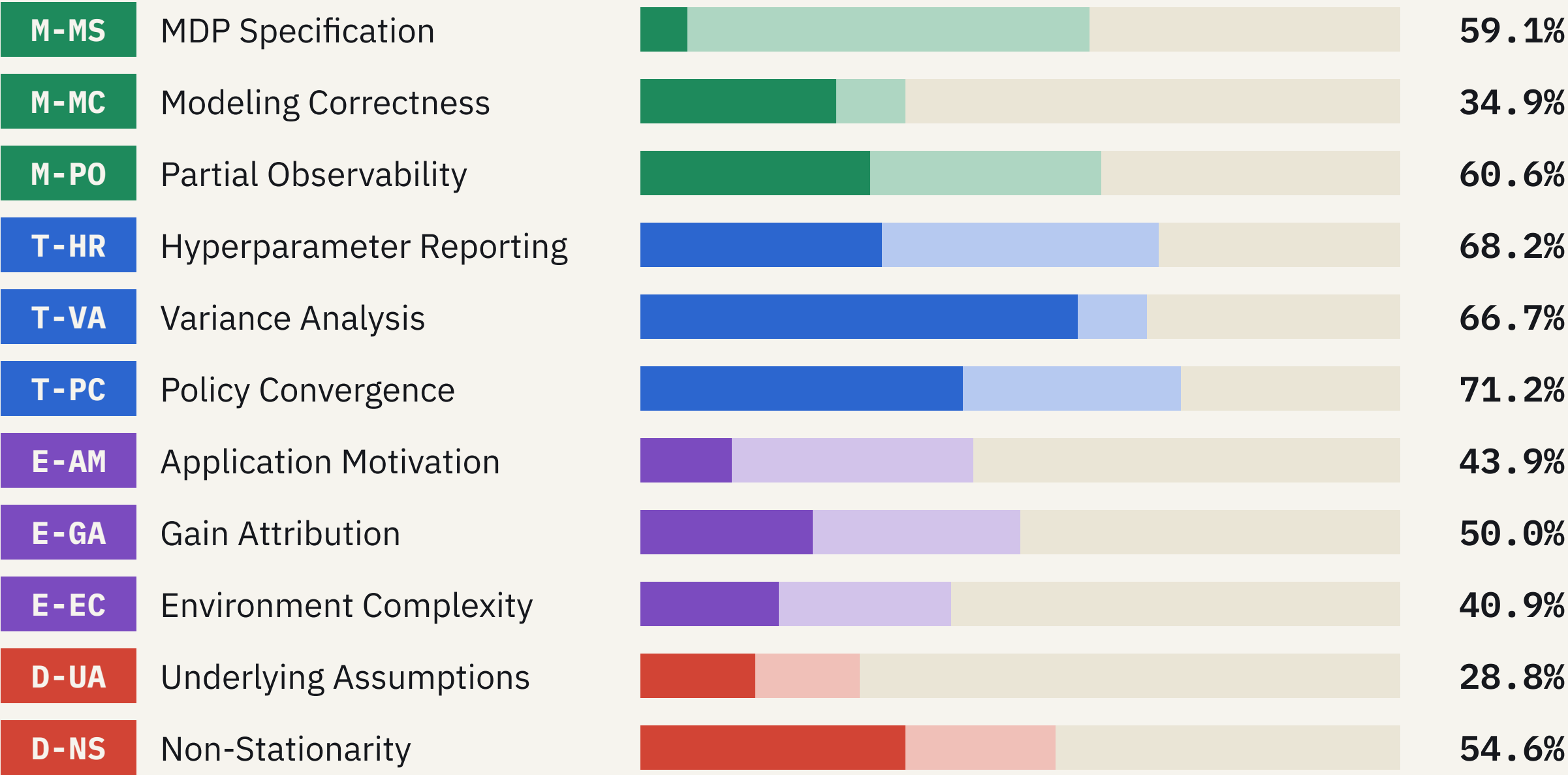
pitfalls identified

66

papers reviewed

≥ 2

pitfalls in every paper



present partially present

Pitfall Prevalence

11

pitfalls identified

66

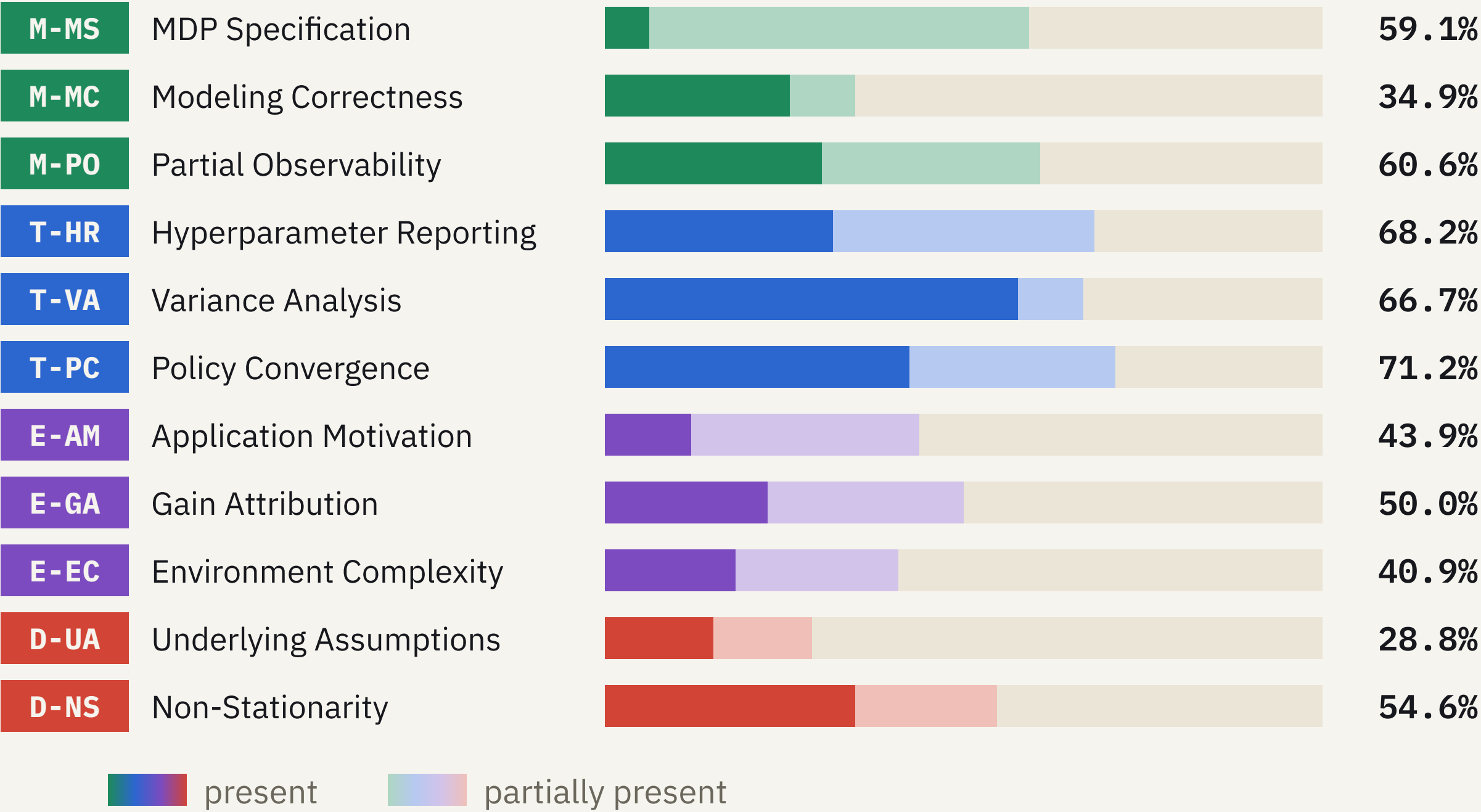
papers reviewed

≥ 2

pitfalls in every paper

5.8

average pitfalls per paper

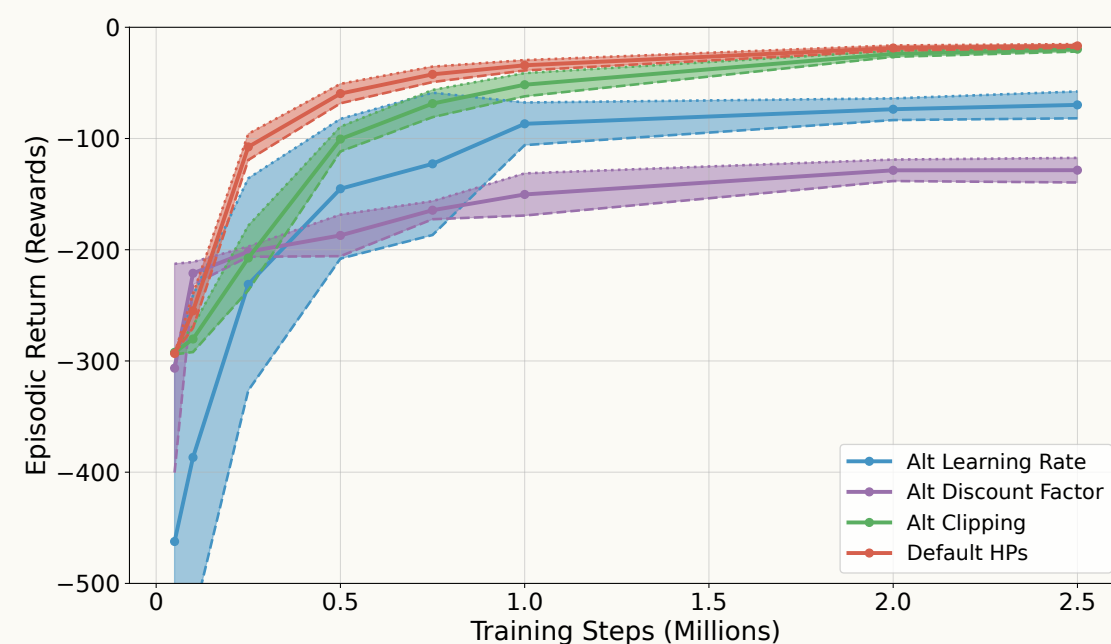


Case Study on Autonomous Cyber Defence

Case Study on Autonomous Cyber Defence

Training Pitfalls

MiniCAGE replicates the CAGE 2 autonomous cyber defense challenge: a blue agent defends a 13 host enterprise network.

**T-HR**

A single hyperparameter can have significant impacts on performance.

T-VA

Single runs can greatly over-estimate performance.

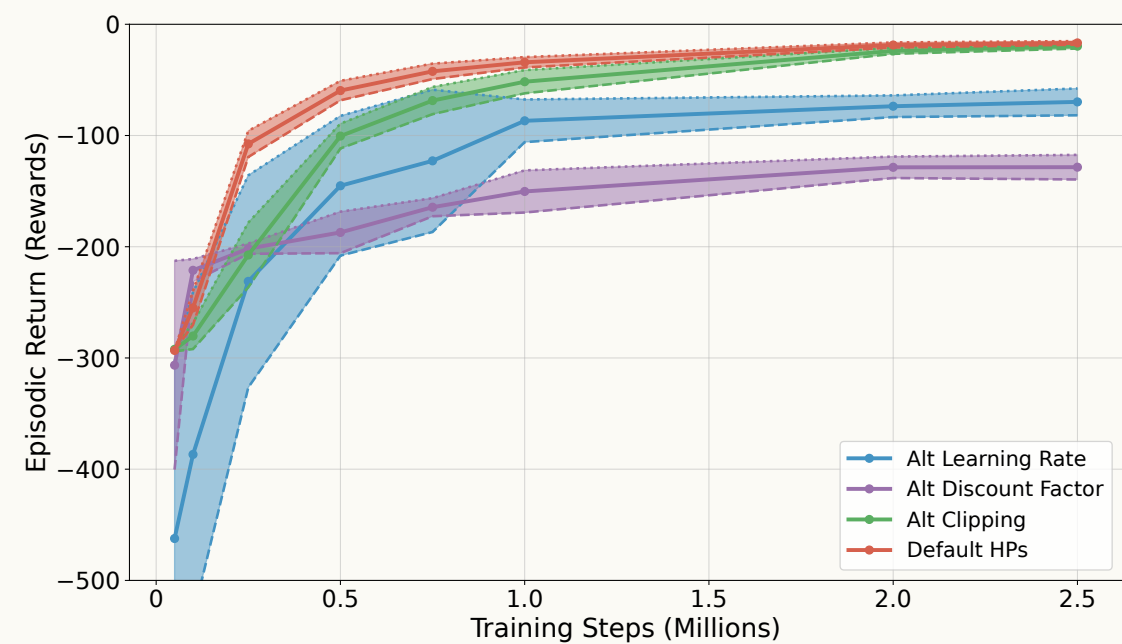
T-PC

Convergence should be demonstrated, e.g., through plots.

Case Study on Autonomous Cyber Defence

Training Pitfalls

MiniCAGE replicates the CAGE 2 autonomous cyber defense challenge: a blue agent defends a 13 host enterprise network.



- T-HR** A single hyperparameter can have significant impacts on performance.
- T-VA** Single runs can greatly over-estimate performance.
- T-PC** Convergence should be demonstrated, e.g., through plots.

Deployment Pitfalls

The blue agent is trained against one attacker strategy or a mix, then evaluated against each.

TRAIN \ EVAL	B-LINE	MEANDER	MIXED
B-line	-19.0	-104.1	-61.4
Meander	-79.0	-61.0	-70.2
Mixed	-40.2	-79.2	-60.2

- B-line** Heads directly to the operational server to compromise it.
- Meander** Compromises every host along the way.
- D-NS** Training on a fixed adversary collapses against unseen attackers.

The Pitfalls of Deep Reinforcement Learning for Cybersecurity

Shae McFadden, Myles Foley, Elizabeth Bates, Ilias Tsingenopoulos,
Sanyam Vyas, Vasilios Mavroudis, Chris Hicks, Fabio Pierazzi



PAPER



WEBSITE

"These are not failures of individual researchers but systematic challenges that emerge when adapting DRL to cybersecurity."

The goal is not to discourage DRL in security. It is to establish methodological standards that distinguish robust, deployable solutions from those that succeed only in idealized settings.