

Towards Scalable Verification of Deep Reinforcement Learning

October 2021

Guy Amir

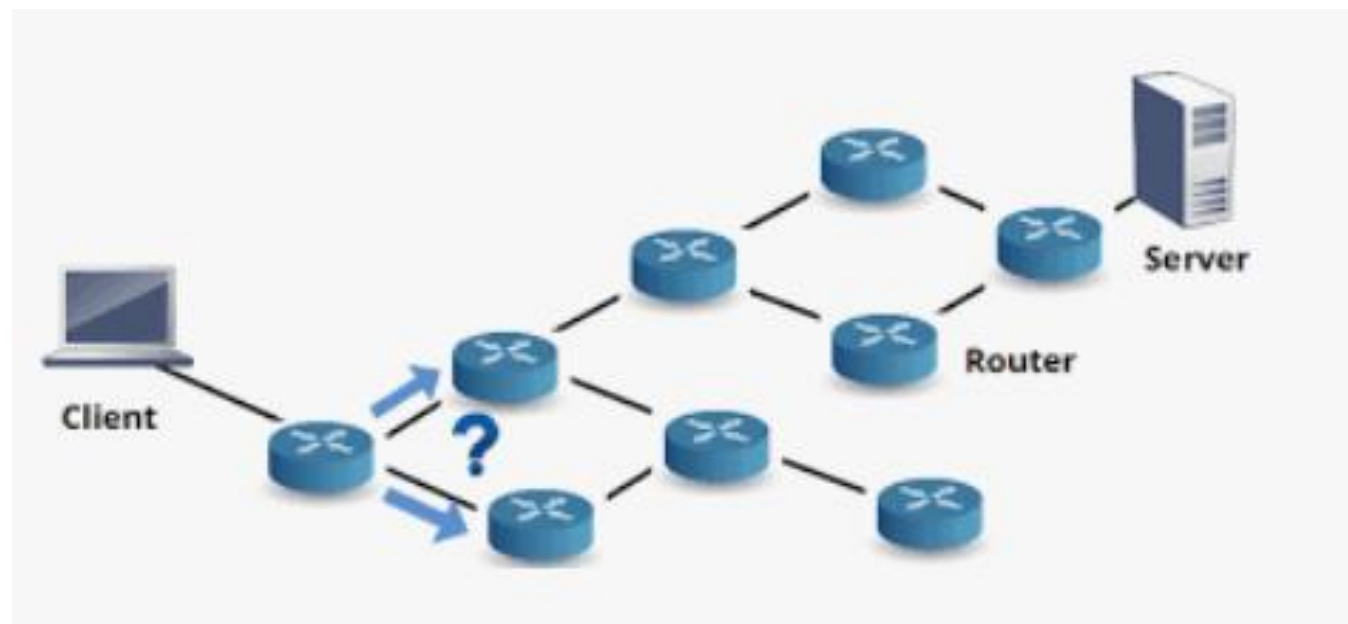
Michael Schapira

Guy Katz



Traditionally

Computer and networked systems are **handcrafted** by domain-specific experts



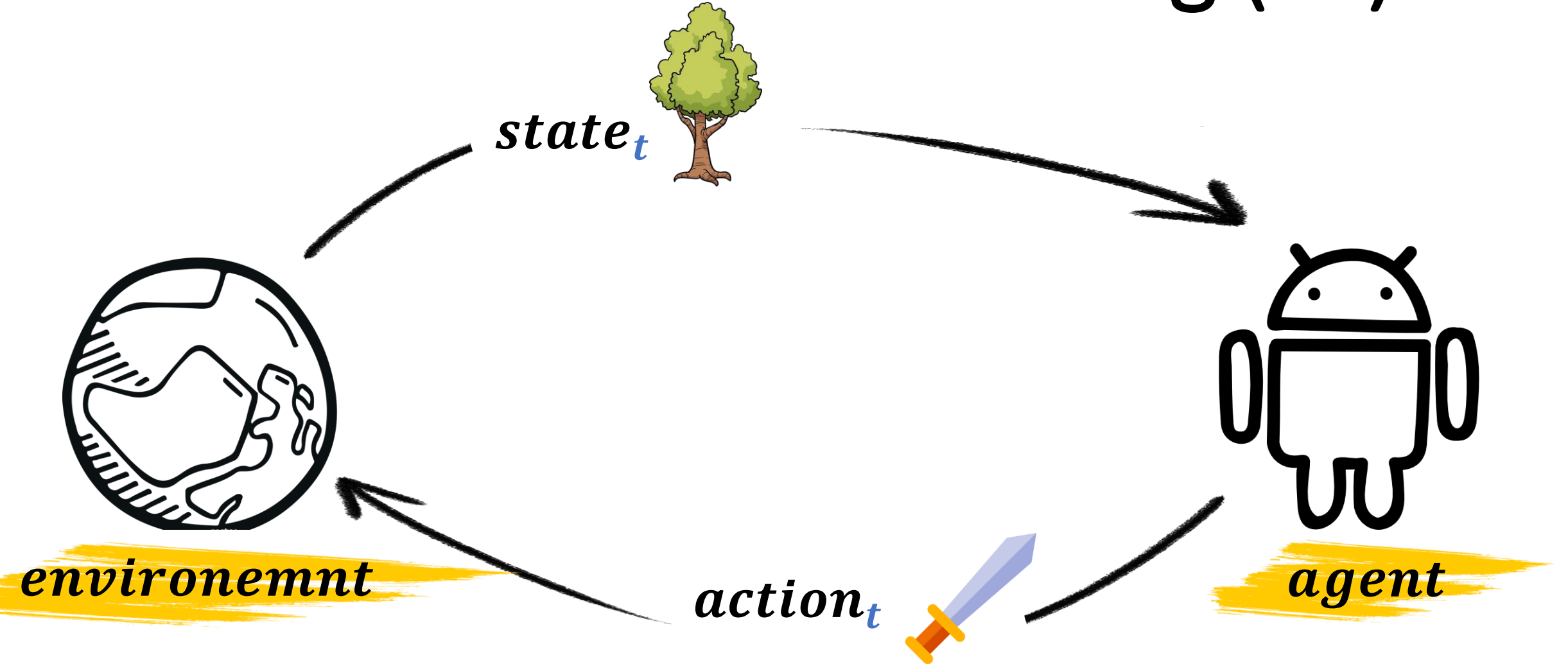
An Emerging Alternative

 Deep **Reinforcement Learning** (DRL) solutions

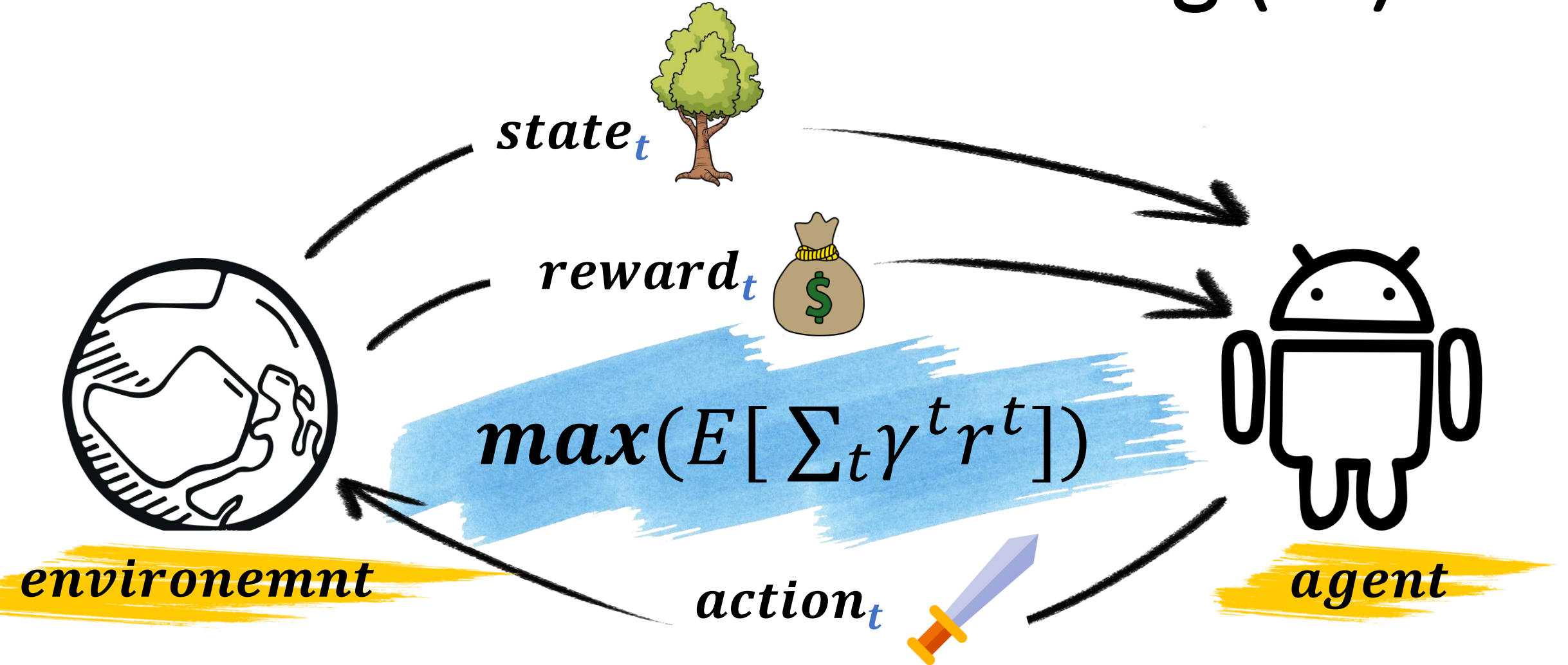
 Outperform the **state-of-the-art** in various contexts

System	Application Domain
Aurora [29]	congestion control
NeuroCuts [40]	packet classification
[51]	SQL optimization
NEO [49]	SQL optimization
DeepRM [44]	resource allocation
[72]	resource allocation
[42]	resource & power management
[36]	compiler phase ordering
[52]	device placement
Placeto [2]	device placement
Decima [48]	spark cluster job scheduling
Pensieve [46]	adaptive video streaming
AuTO [11]	traffic optimizations

Reinforcement Learning (RL)

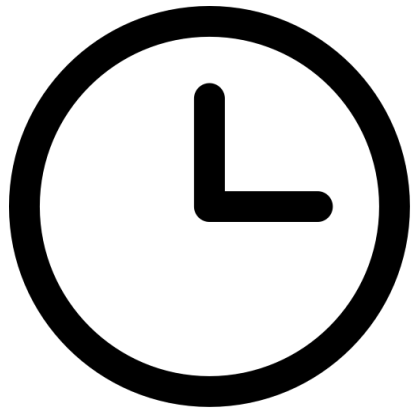


Reinforcement Learning (RL)

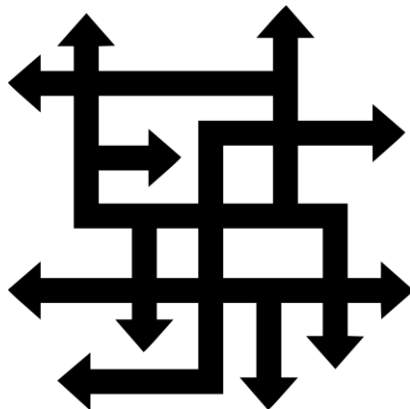


Reinforcement Learning (RL)

Infinite
Runs



Complex
Policies



Communication
Domain



But...



*How do we know a Deep Neural Network trained via **Reinforcement Learning** is safe?*

“Testing shows the presence, not the absence of bugs”

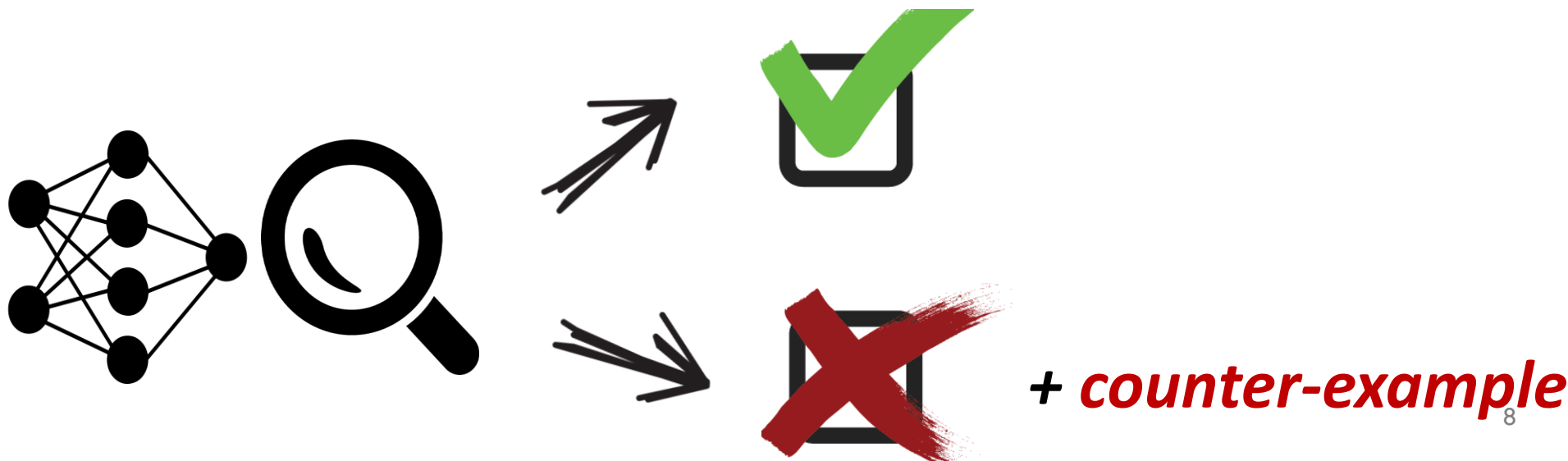
Dijkstra, 1969



Challenge: These “black boxes” need to be **formally verified** for correct behavior

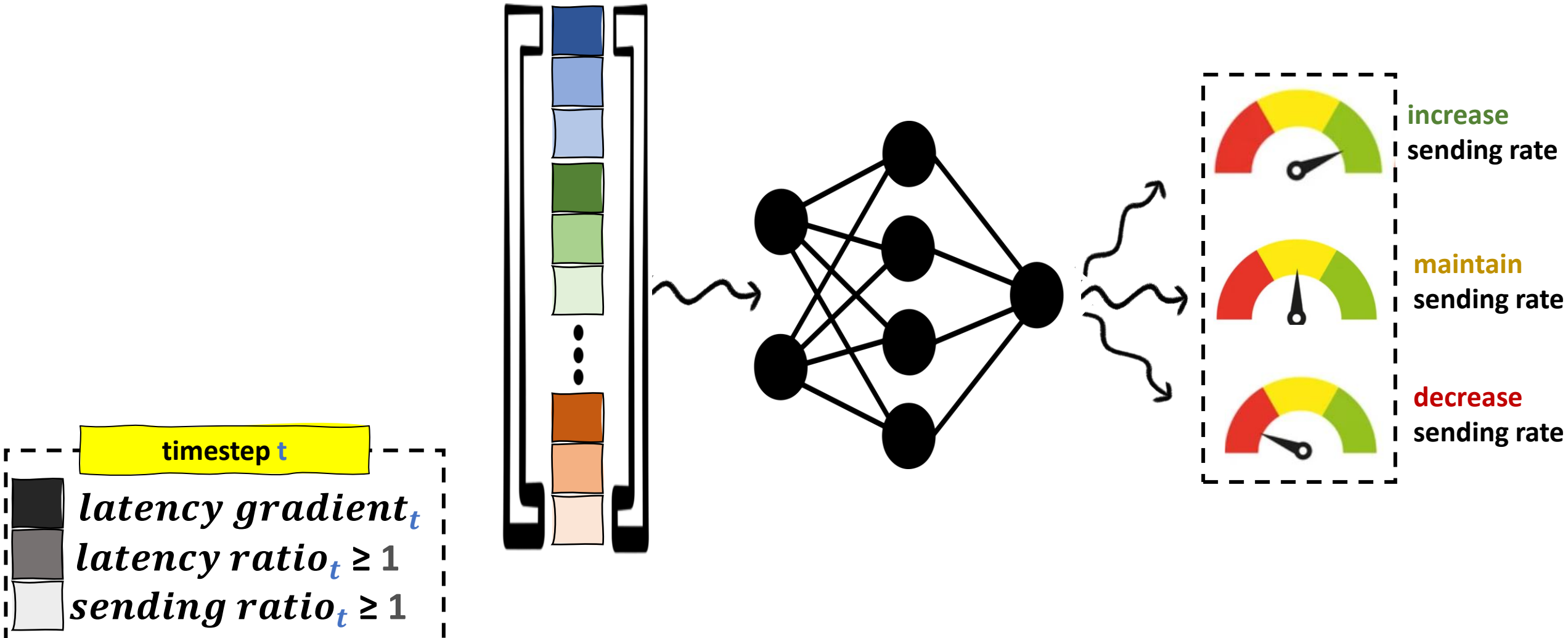
Our approach: Formal Verification!

*Provably **guarantee** that a learned policy meets our requirements, or identify concrete **violations** (bugs)*

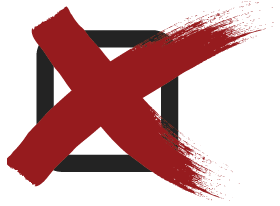


Example: The **Aurora** Congestion Controller

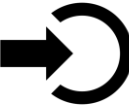
[Jay, Rotman, et al., ICML 2019]



Aurora **Safety** Properties



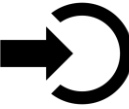
Safety - “Something **bad** never happens”
(**finite**-long violations)

 Network Conditions	 Wanted Output
poor 	next-step decrease 
excellent 	next-step increase 

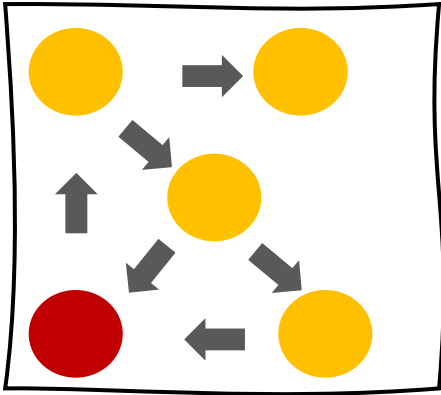
Aurora Liveness Properties



Liveness - “Something **good** eventually happens”
(infinite-long violations)

 Network Conditions	 Wanted Output
poor 	eventual decrease 
excellent 	eventual increase 

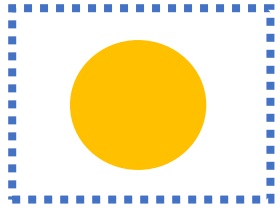
Our Verification Strategy



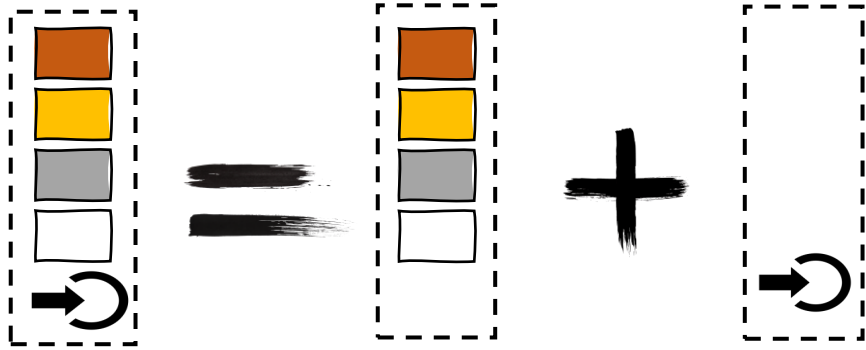
Defining a **state graph**
& **transition function**

[Eliyahu-Kazak-Katz-Schapira, SIGCOMM 2021]

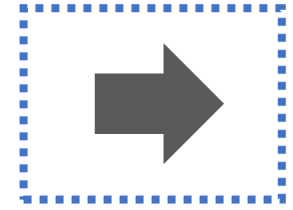
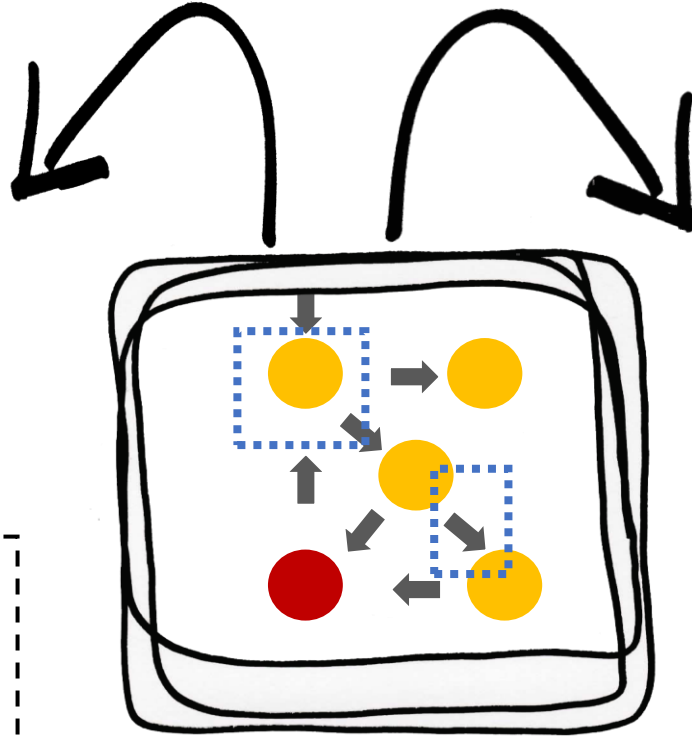
Transition System Graph



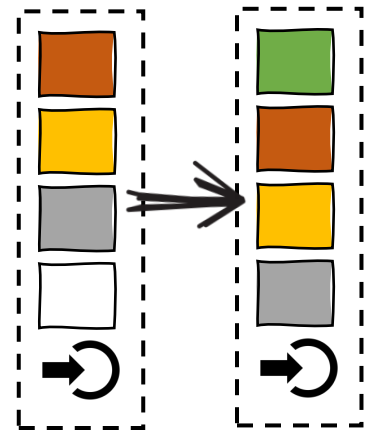
Defining a $state_t$



$state_t$ $input_t$ $output_t$

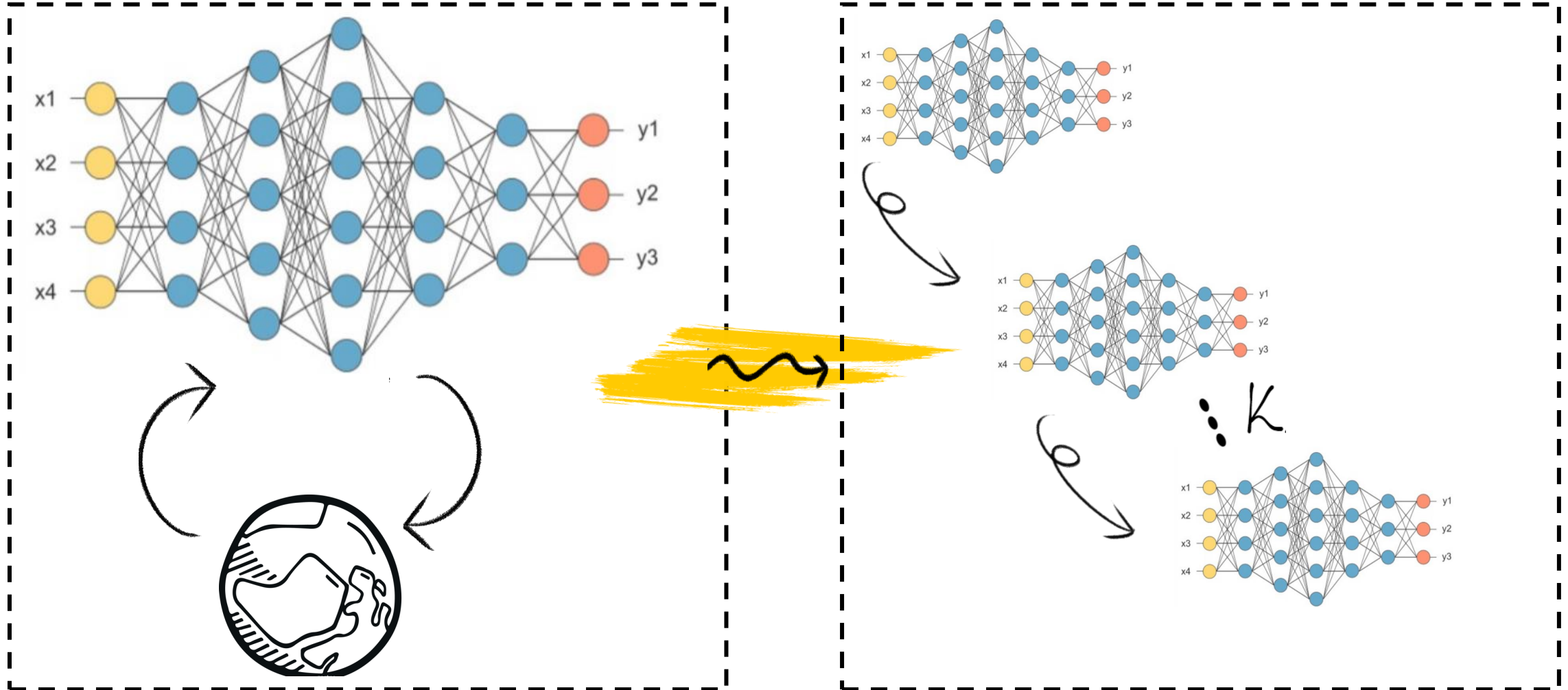


Defining a $transition_{t,t'}$

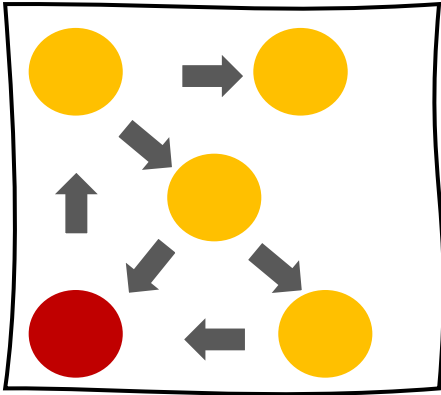


$state_t$ $state_{t'}$

Encoding Multiple Transitions

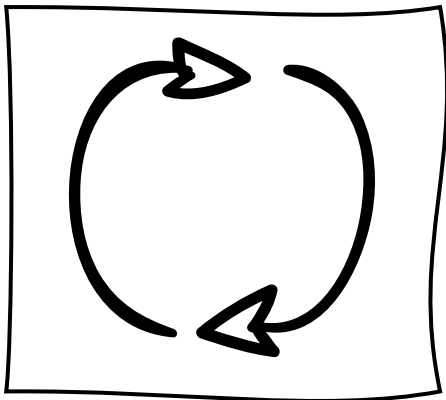


Our Verification Strategy



Defining a **state graph**
& **transition function**

[Eliyahu-Kazak-Katz-Schapira, SIGCOMM 2021]



Running a **portfolio approach** for checking
k-long **violations** or **k**-long **provable runs**

[Amir-Schapira-Katz, FMCAD 2021]

Bounded Model Checking (BMC)

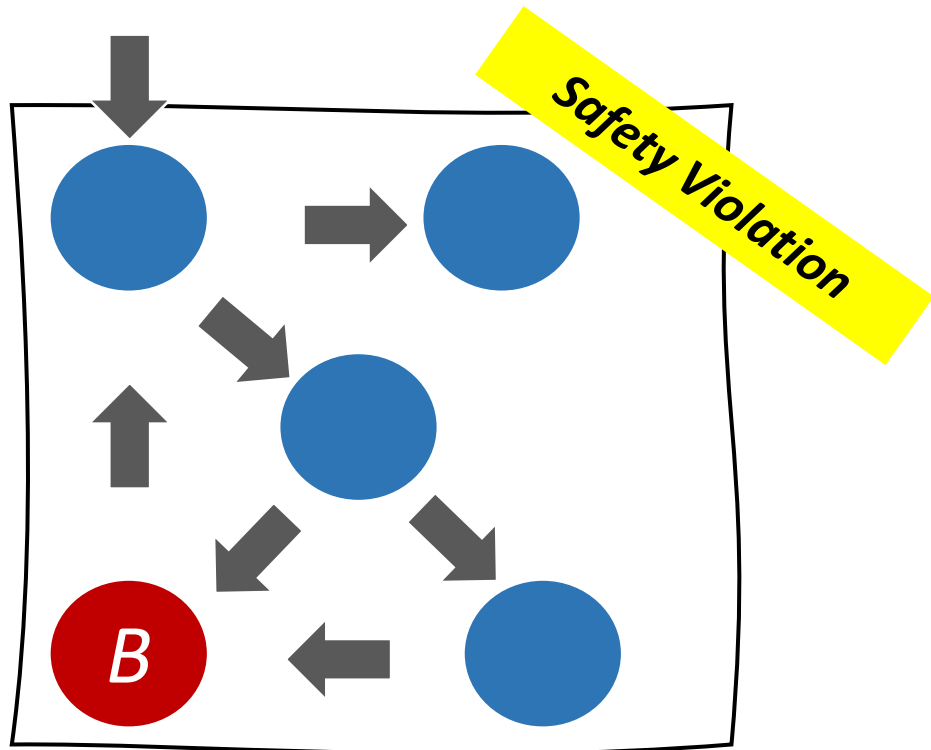
Bounded Model Checking

A method for checking **violations** of properties, for a given number of **k steps**

Bounded Model Checking (BMC)

Bounded Model Checking

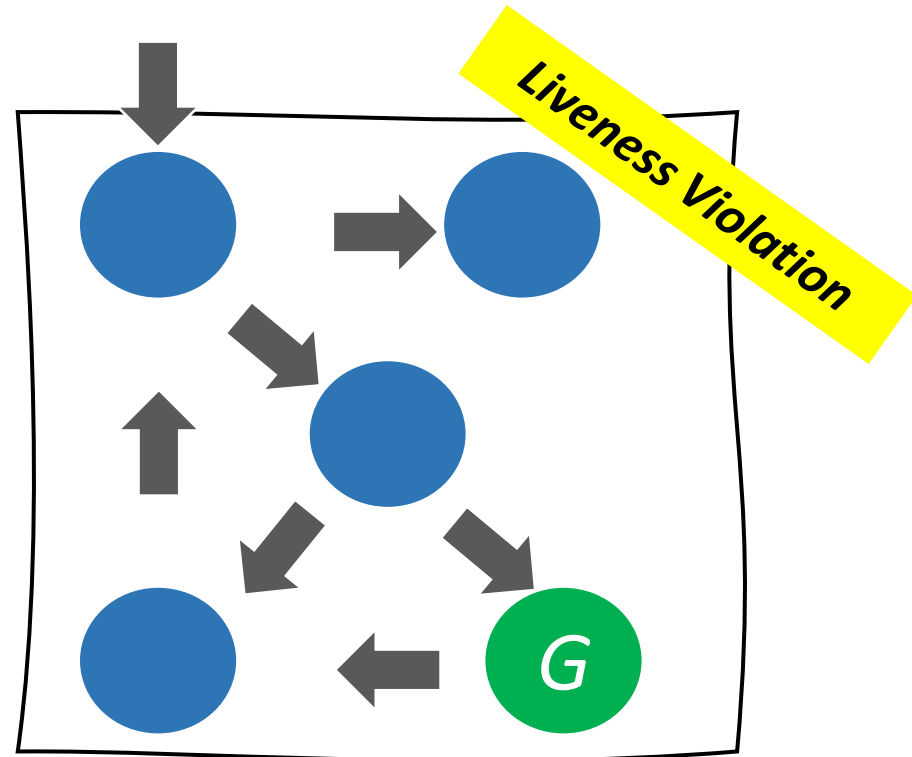
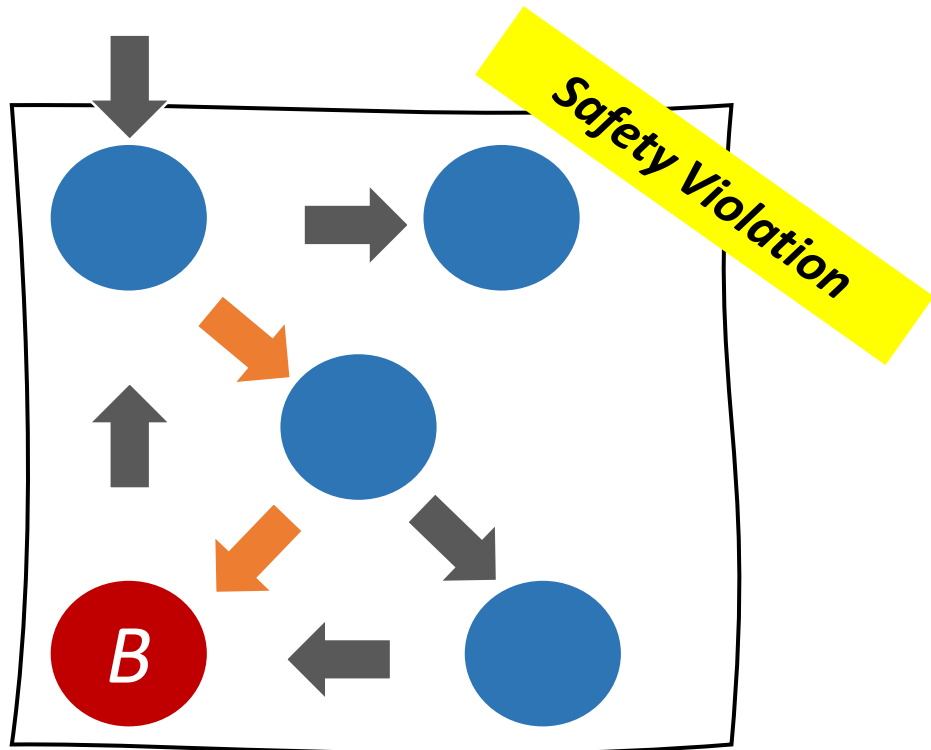
A method for checking **violations** of properties, for a given number of **k steps**



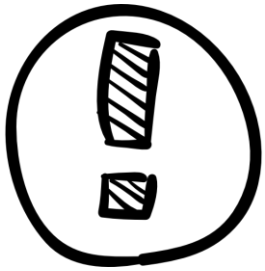
Bounded Model Checking (BMC)

Bounded Model Checking

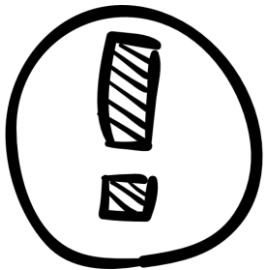
A method for checking **violations** of properties, for a given number of **k steps**



BMC Setbacks

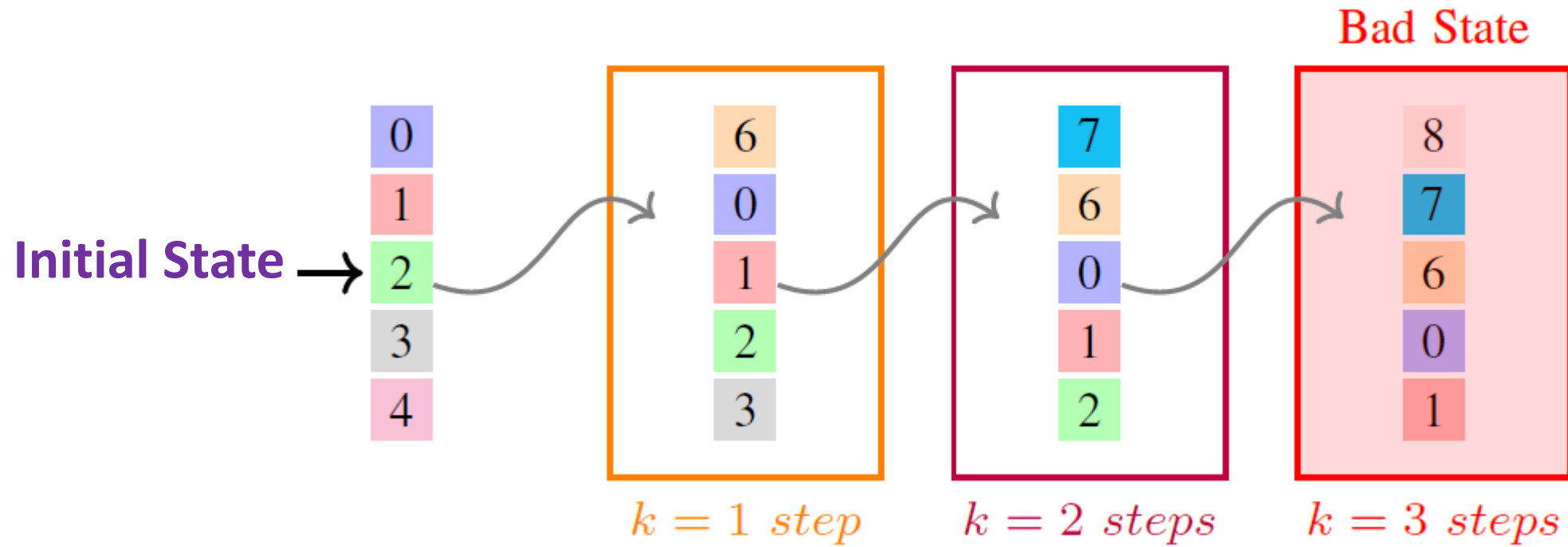


*We **can't** prove that any **properties** hold*

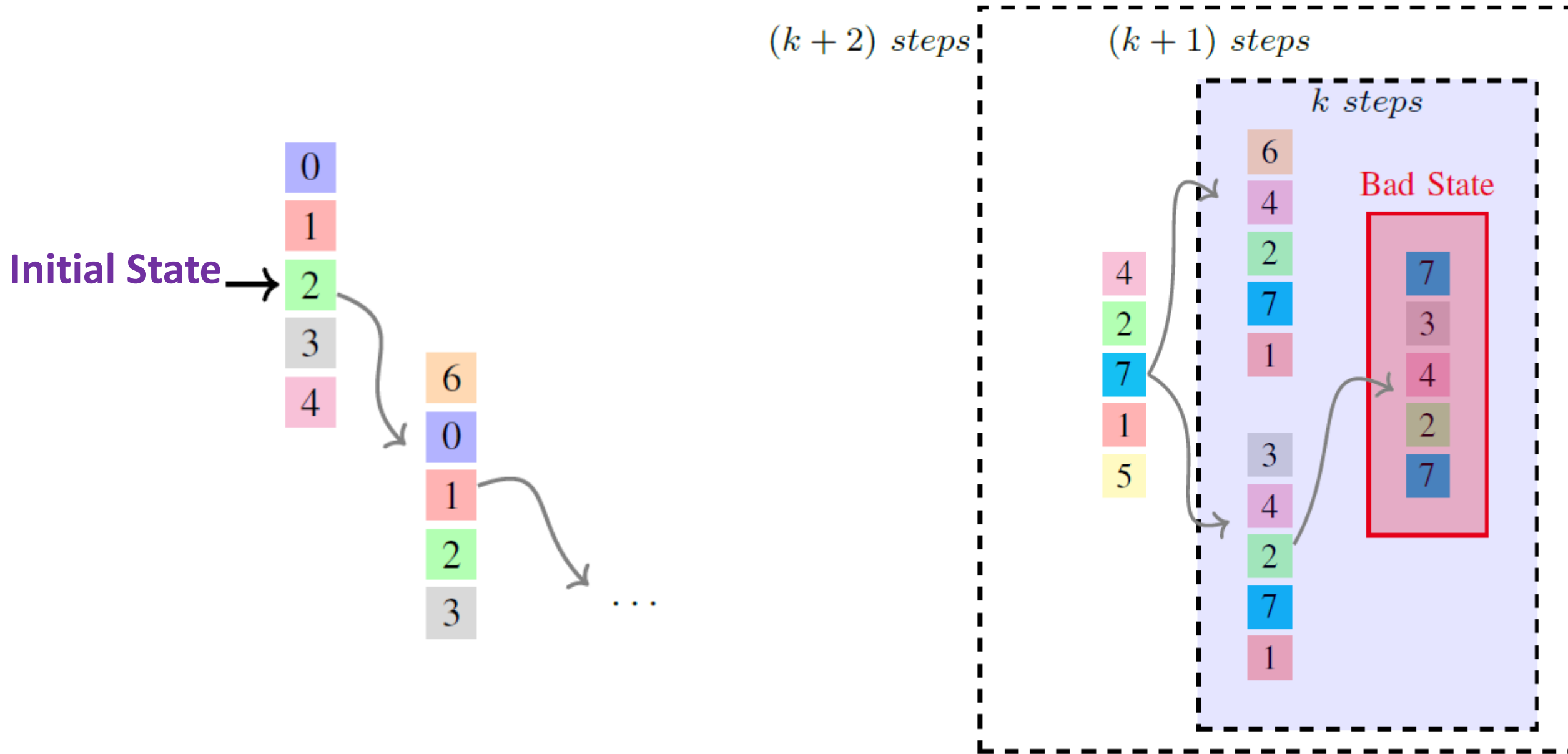


*We **can't** analyze **complex properties***

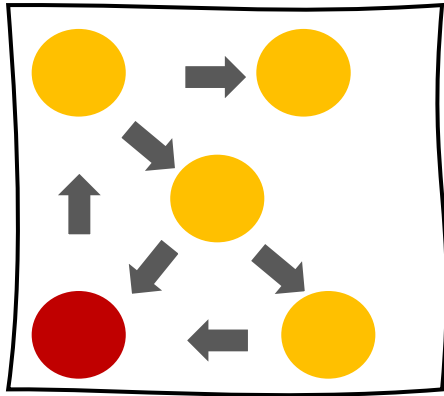
BMC



K-Induction

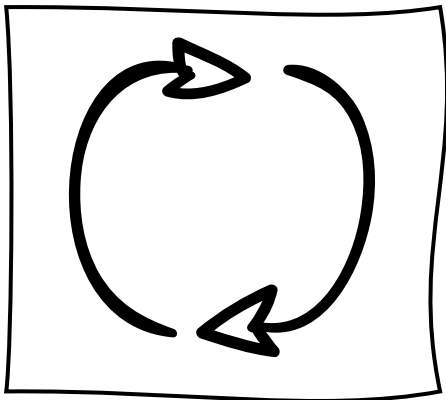


Our Verification Strategy



Defining a **state graph**
& **transition function**

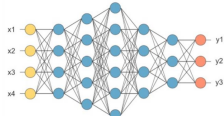
[Eliyahu-Kazak-Katz-Schapira, SIGCOMM 2021]



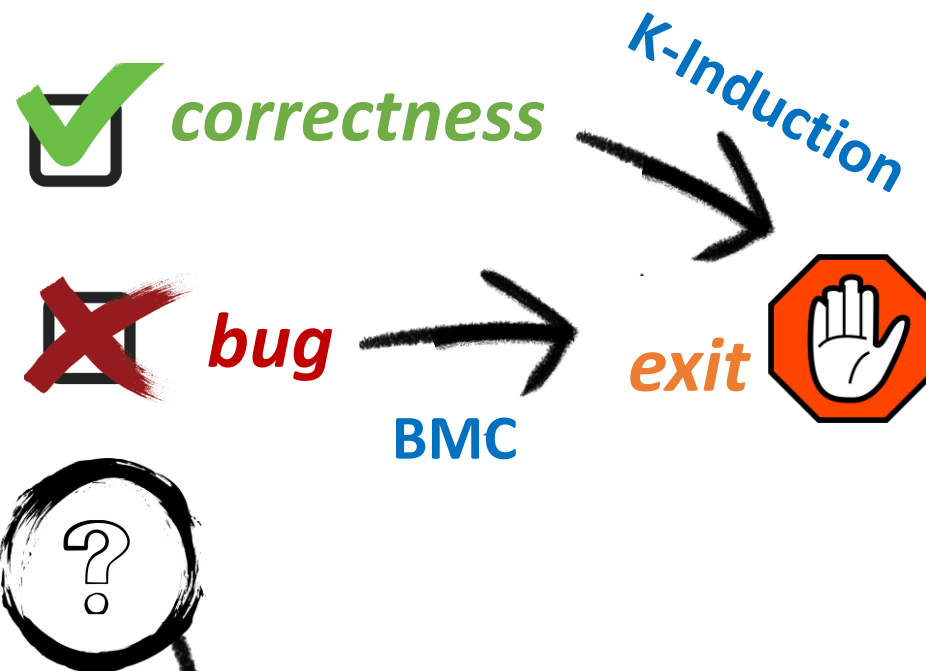
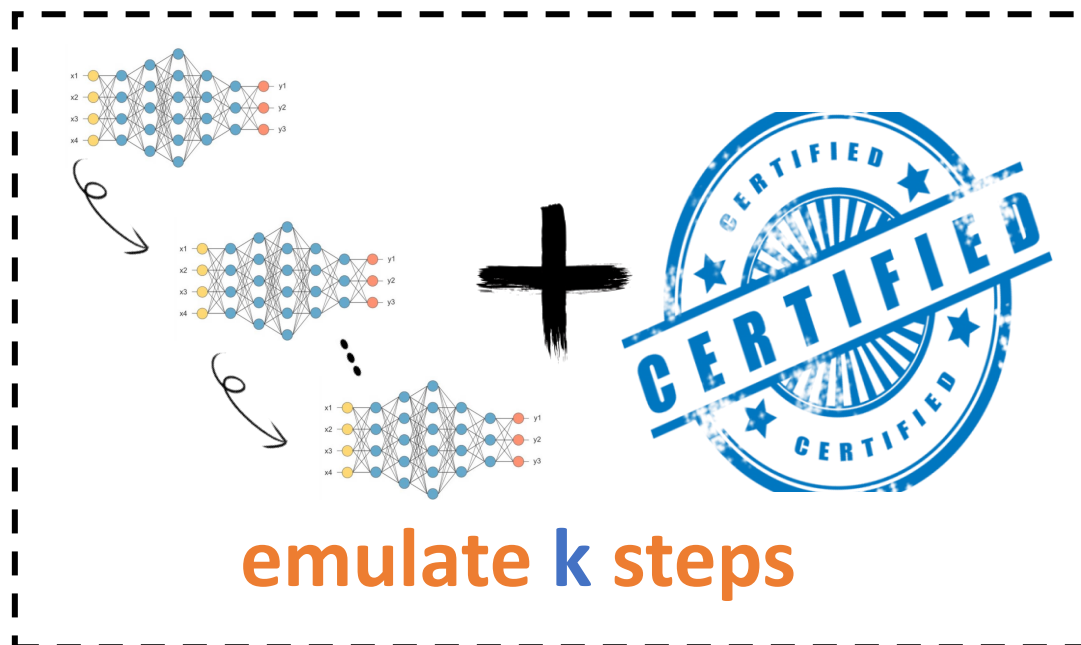
Running a **portfolio approach** for checking
k-long **violations** or **k**-long **provable runs**

[Amir-Schapira-Katz, FMCAD 2021]

Portfolio Approach

input:  + **property**

initialization: **k=1** step

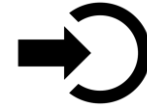


increment: **k ++**

Aurora Properties



Network Conditions



Desired Output

Property 1

excellent 😊

eventual change 👍👎



Property 2

excellent 😊

eventual increase 👍



Property 3

poor 😞

next-step decrease 👎



Property 4

poor 😞

eventual decrease 👎



WhiRL 2.0 - Techniques

1

K-Induction

2

Invariant Inference

3

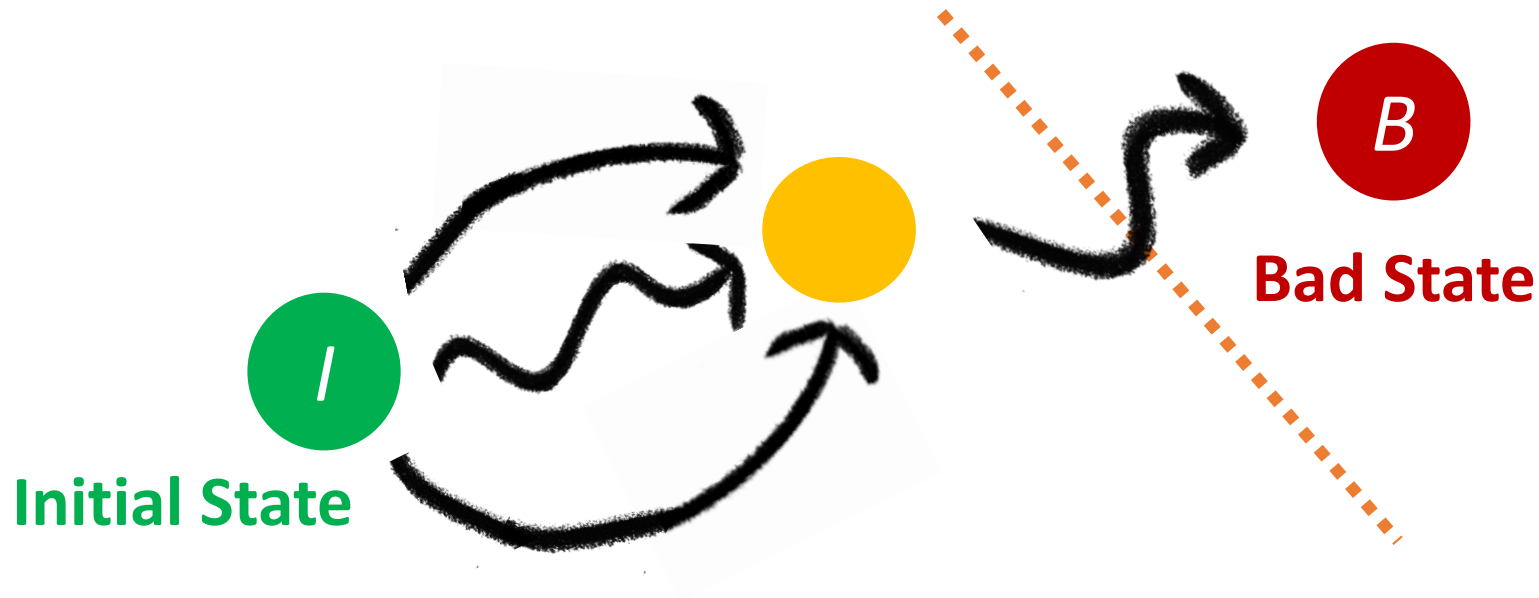
Abstraction

Invariant Inference

Invariant

A **partition of the state space** S into two disjoint sets S_1 and S_2 such that:

$$s_1 \in S_1 \wedge s_2 \in S_2 \rightarrow (s_1, s_2) \notin \text{transition } T$$



Invariant Inference



Templates:

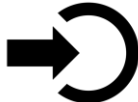
- ✓ use *monotonicity* of properties
- ✓ fix *inputs* or *outputs*
- ✓ conduct a *binary search* on the non-fixed variables
- ✓ *dynamic*: user-chosen values



Strategy: search for the “*2nd best*” behavior

Invariant Inference - Aurora

- Execution of *invariance inference algorithms* based on the *templates*
- For example, an invariant is found, *based on the following violated property*

	 Network Conditions	 Wanted Output	 Property Holds?
Property 3	poor 	next-step decrease 	 k=1

Invariant Inference - Aurora



Originally, “*poor*” network conditions:

$$2 \leq \textit{sending_ratio}_t \quad \text{😞}$$



We can search for the *worst-case sending ratio* for the output to decrease:

$$\textit{output}_t < 0 \quad \text{👎}$$

Invariant Inference - Aurora

 Initialization: $0 \leq output$, $2 \leq sending_ratio_t \leq M_{user}$

 Iterate:

 binary-search the $sending_ratio_t$ lower bound

 call a *verifier* on the middle point

 update $sending_ratio_t$ accordingly

 Return: lower bound on worst case $sending_ratio_t$

Invariant Inference - Aurora



$sending_ratio_t$ lower bound for **SAT**

→ **2**



$sending_ratio_t$ lower bound for **UNSAT**

→ **M**

$verifier\{sending_ratio_t \in [2, M]\} \rightarrow \text{SAT}$



2



$\frac{1}{2}(2+M)$



M

$verifier\{sending_ratio_t \in [\frac{1}{2}(2+M), M]\} \rightarrow \text{SAT}$

Invariant Inference - Aurora



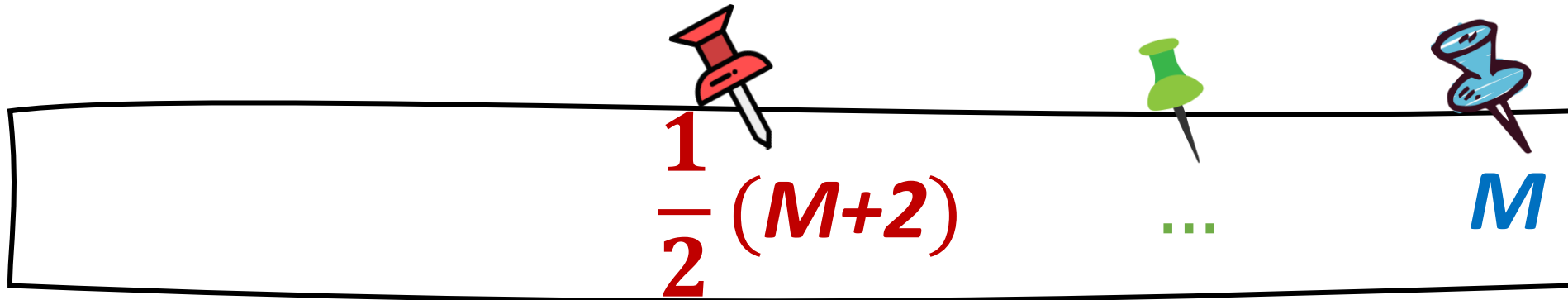
sending_ratio_t lower bound for **SAT**

→ **2**



sending_ratio_t lower bound for **UNSAT**

→ **M**



$$\text{verifier}\{ \textit{sending_ratio}_t \in [\frac{1}{2}(\frac{1}{2}(M+2) + M), M] \} \rightarrow \text{UNSAT}$$

Invariant Inference - Aurora



sending_ratio_t lower bound for **SAT**

→ **2**



sending_ratio_t lower bound for **UNSAT**

→ **M**



$\frac{1}{2}(M+2)$

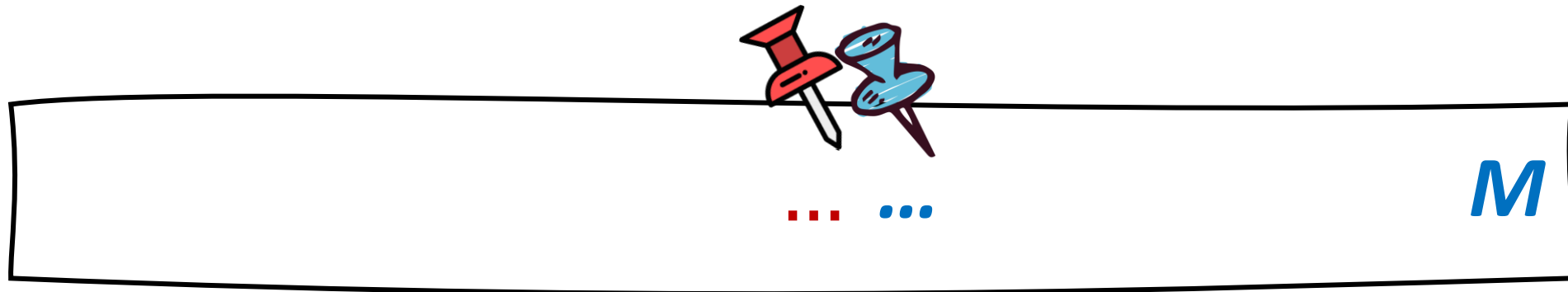


...

M

Invariant Inference - Aurora

after $\log(M)$ iterations:



return:  = lower bound on worst-case $\textit{sending_ratio}_t$

Techniques

1

K-Induction

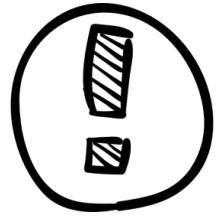
2

Invariant Inference

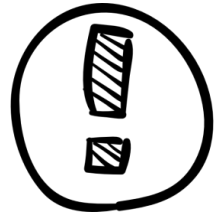
3

Abstraction

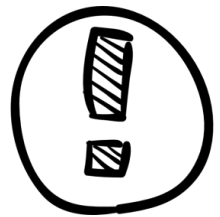
See paper for...



Abstraction techniques for **generalization**



Methods for identifying **undesirable policies**



Modules for improving **interpretability**

Summary



A (first?) method for **proving properties** of RL-driven systems



Automatic **invariant inference** of “2nd best” properties, in chosen scenarios



Explainability and **interpretability** of bad policies

Future Steps



Improve **scalability**



Focus on **generalization**

Questions

