

DeepLTL: Learning to Efficiently Satisfy Complex LTL Specifications for Multi-Task RL

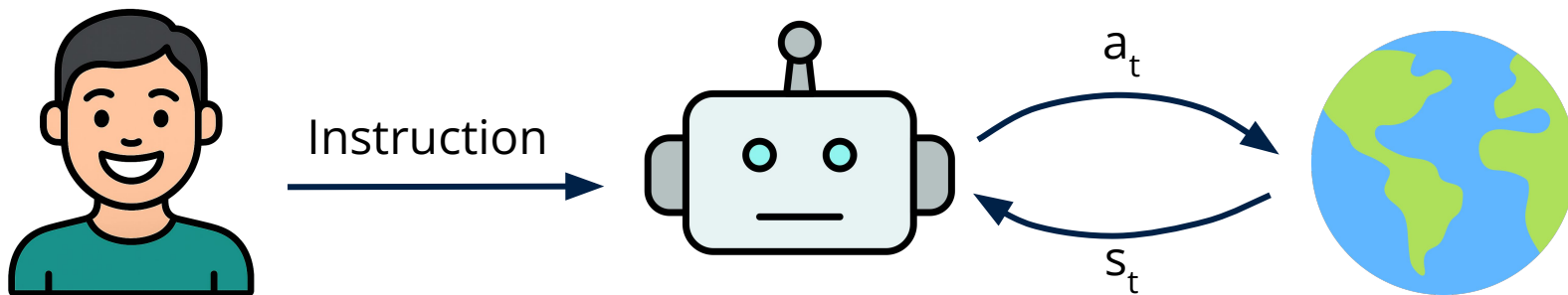
Mathias Jackermeier, Alessandro Abate

ICLR 2025, Singapore

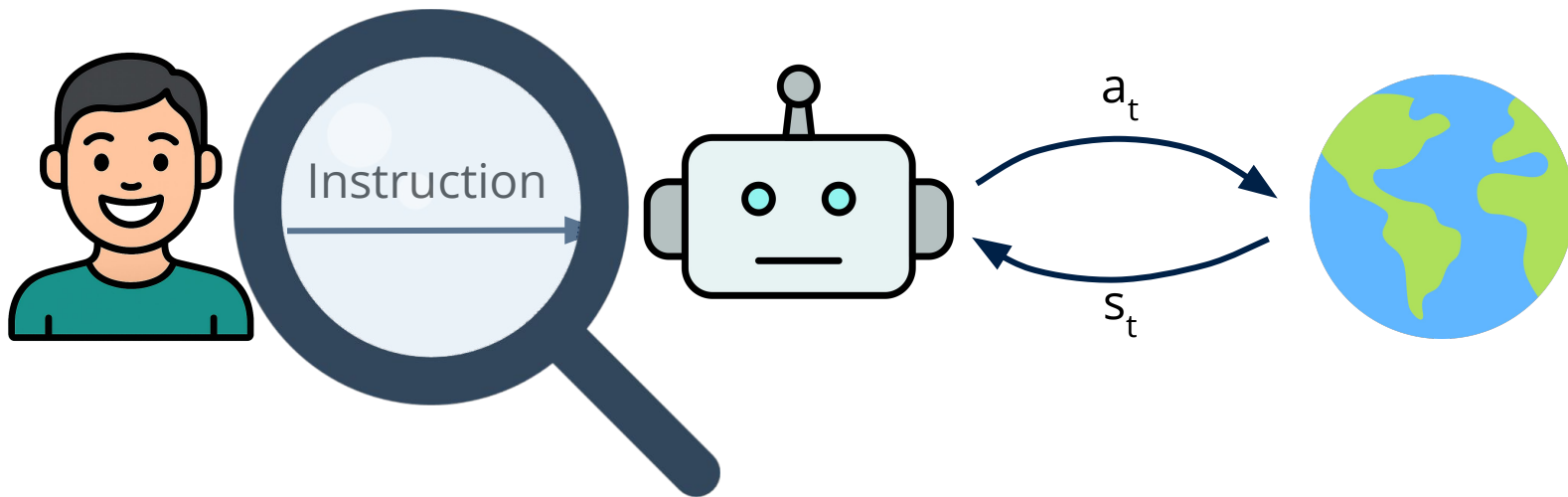


**Engineering and
Physical Sciences
Research Council**

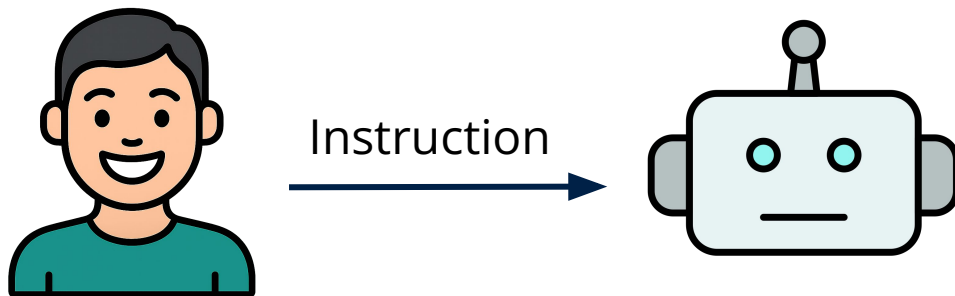
Instruction-following RL agents



Instruction-following RL agents



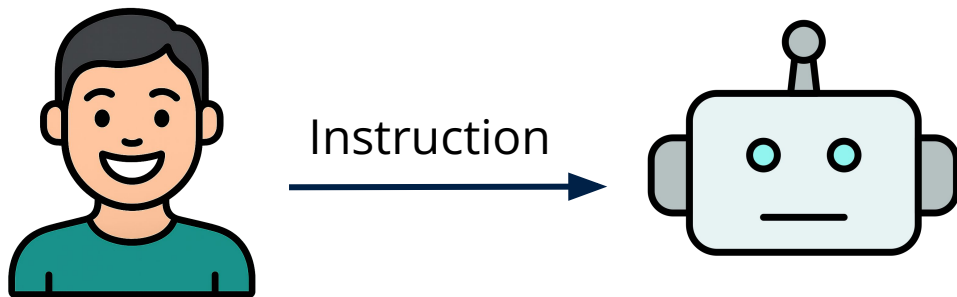
Instruction-following RL agents



Natural language:

- ✓ Intuitive
- ✗ Ambiguous
- ✗ Difficult to assess

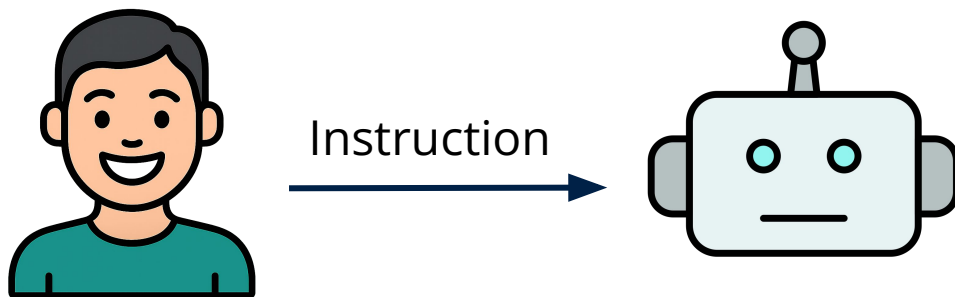
Instruction-following RL agents



Formal specifications:

- ✓ Precise
- ✓ Easy to verify
- ✓ Explicit structure
- ✗ Difficult to formulate (?)

Instruction-following RL agents



Formal specifications:

- ✓ Precise
- ✓ Easy to verify
- ✓ Explicit structure
- ✗ Difficult to formulate (?)

→ Well suited when correctness is crucial, e.g. safety-critical settings

Linear temporal logic (LTL)

Example specification:

$$(\neg \text{yellow} \text{ U } \text{purple}) \wedge G (\text{green} \Rightarrow F \text{blue})$$

Linear temporal logic (LTL)

Example specification:

$$\underbrace{(\neg \text{yellow} \text{ U } \text{purple})}_{\text{Go to the purple zone while avoiding the yellow region,}} \wedge G (\text{green} \Rightarrow F \text{blue})$$

“Go to the purple zone while avoiding the yellow region,

Linear temporal logic (LTL)

Example specification:

$$\mathbf{H} (\neg \text{yellow} \mathbf{U} \text{purple}) \wedge \mathbf{G} (\text{green} \Rightarrow \mathbf{F} \text{blue})$$

“Go to the purple zone while avoiding the yellow region, and

Linear temporal logic (LTL)

Example specification:

$$\boxed{H} (\neg \text{yellow} \cup \text{purple}) \wedge G (\text{green} \Rightarrow F \text{blue})$$

“Go to the purple zone while avoiding the yellow region, and always,

Linear temporal logic (LTL)

Example specification:

$$(\neg \text{yellow} \text{ U } \text{purple}) \wedge \text{G} \underbrace{(\text{green} \Rightarrow \text{F blue})}$$

“Go to the purple zone while avoiding the yellow region, and always, if you visit green you eventually have to go to blue.”

Linear temporal logic (LTL)

Example specification:

$$(\neg \text{yellow} \text{ U } \text{purple}) \wedge G (\text{green} \Rightarrow F \text{blue})$$

“Go to the purple zone while avoiding the yellow region, and always, if you visit green you eventually have to go to blue.”

Linear temporal logic (LTL)

Example specification:

$$(\neg \text{yellow} \text{ U } \text{purple}) \wedge G (\text{green} \Rightarrow F \text{blue})$$

“Go to the purple zone while avoiding the yellow region, and always, if you visit green you eventually have to go to blue.”



Compositional



Temporally
extended



Infinite
horizon



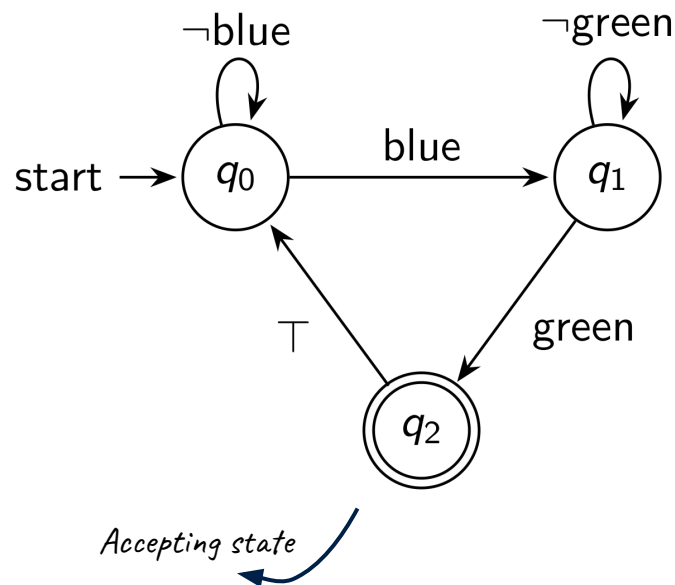
Safety
constraints

How can we train a **multi-task** policy to **zero-shot** execute **arbitrary** LTL specifications?

From LTL specifications to automata

Any LTL specification can be converted to an equivalent (Büchi) **automaton**:

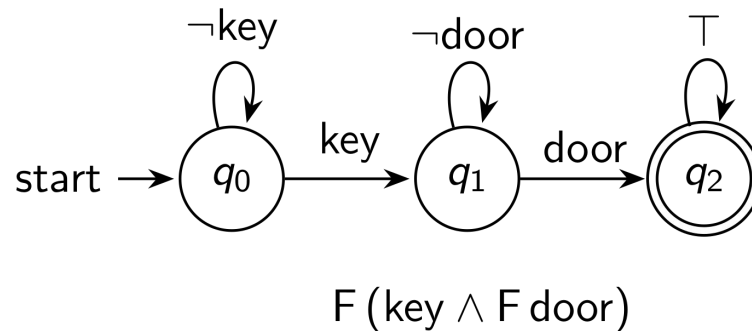
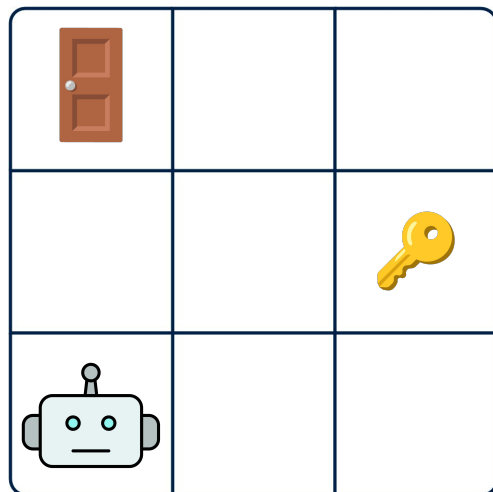
$G F \text{blue} \wedge G F \text{green}$ 



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

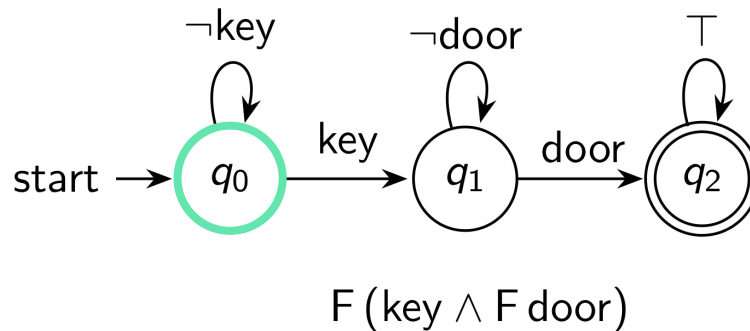
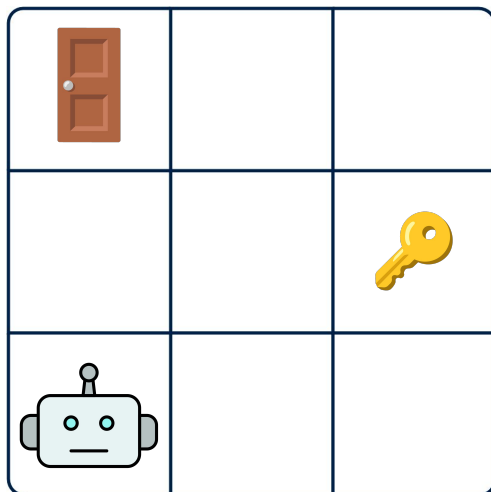
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

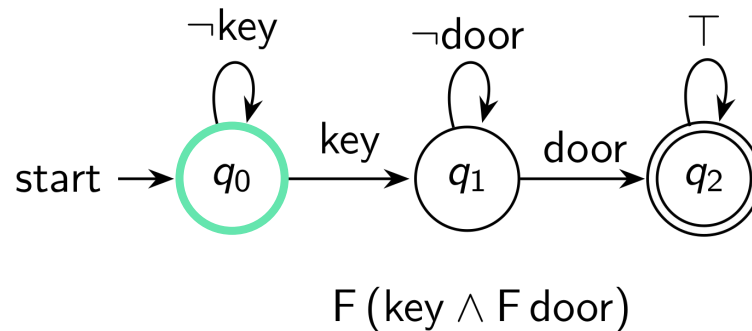
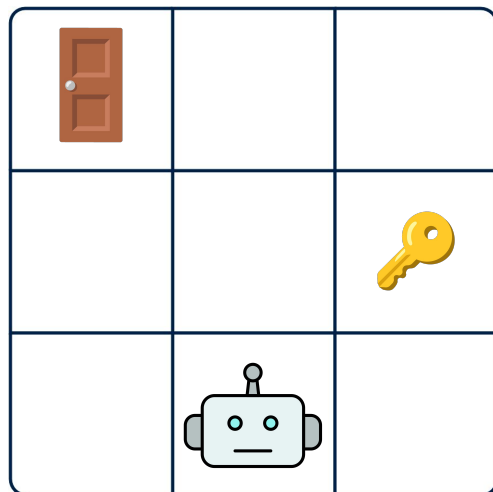
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

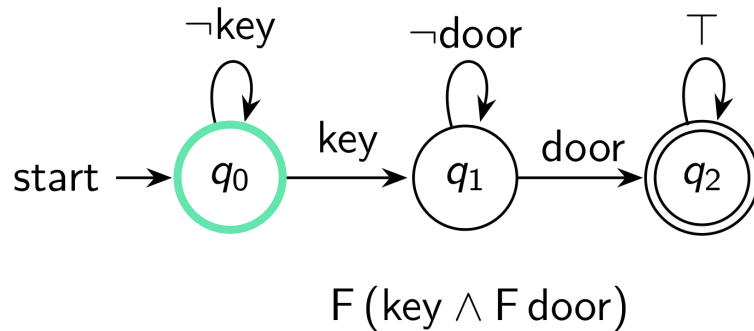
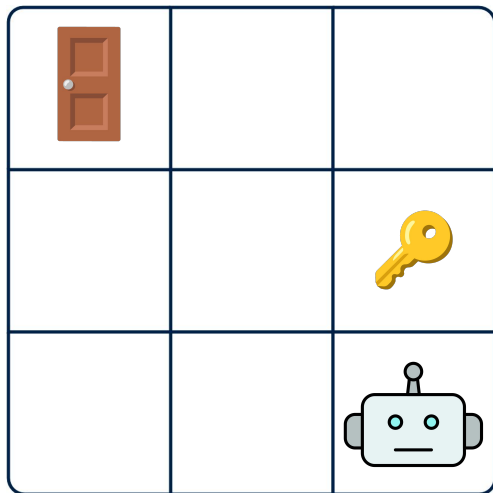
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

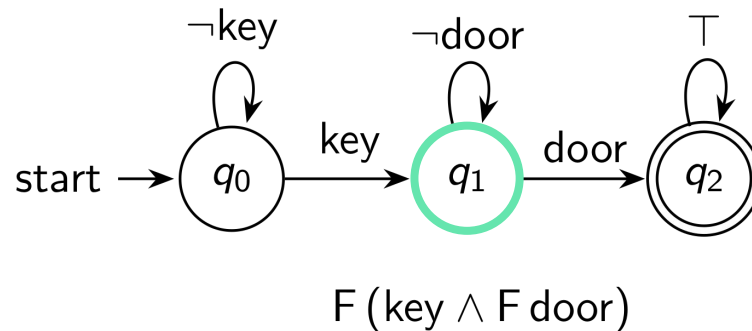
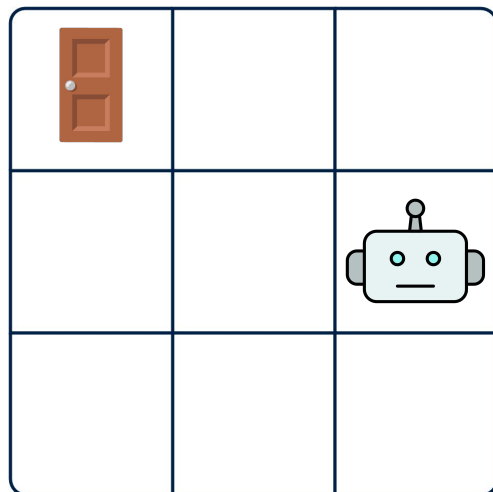
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

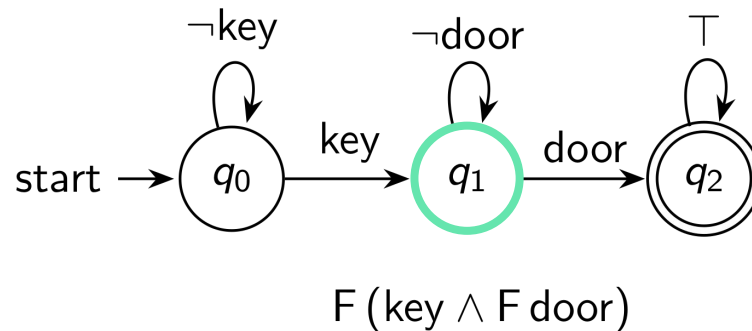
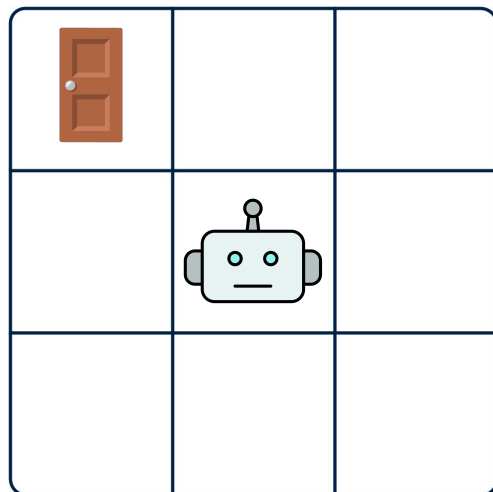
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

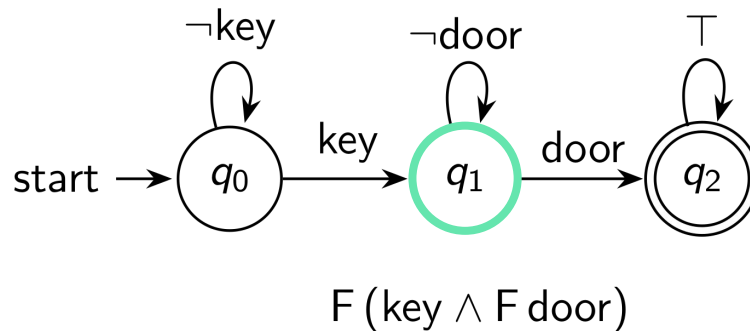
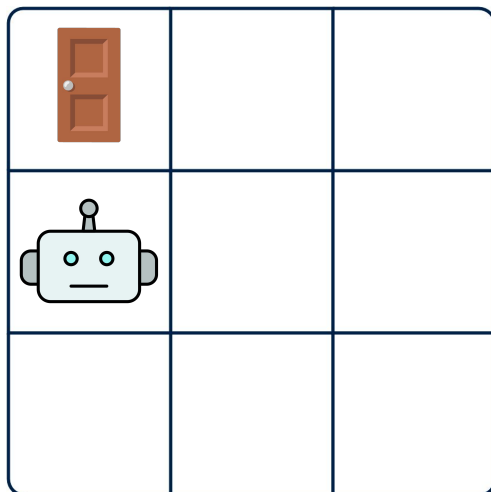
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

Keeping track of the automaton state allows us to learn a **Markovian** policy

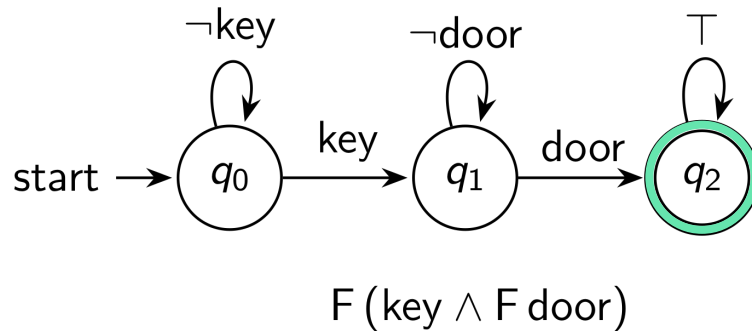
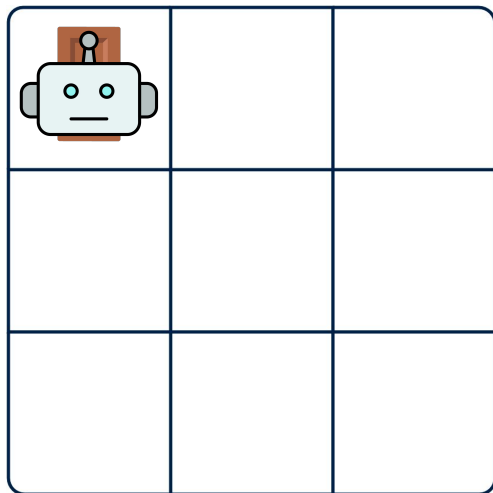
$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



Product MDP

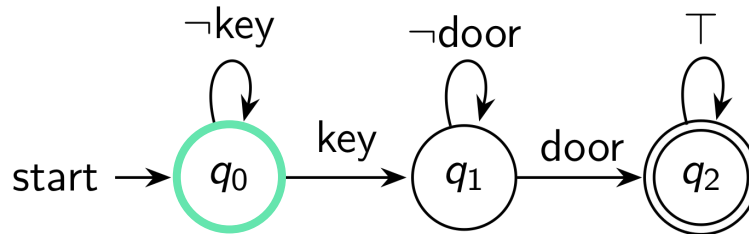
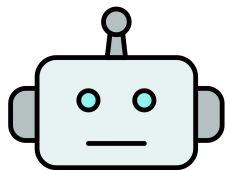
Keeping track of the automaton state allows us to learn a **Markovian** policy

$$\pi : \mathcal{S} \times \mathcal{Q} \rightarrow \Delta(\mathcal{A})$$



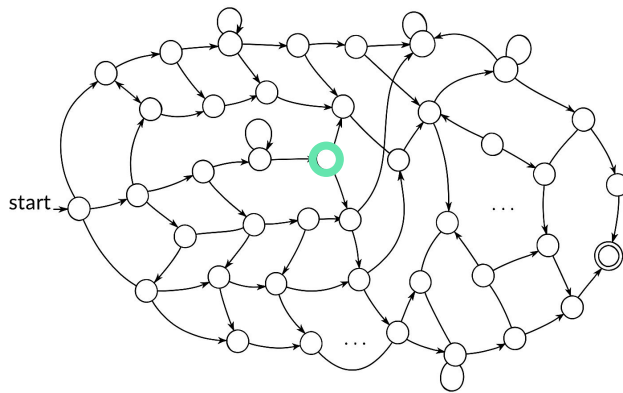
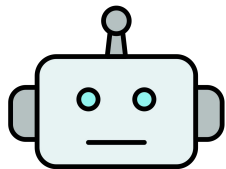
From single-task to multi-task

In a multi-task setting, we do not know the automaton beforehand



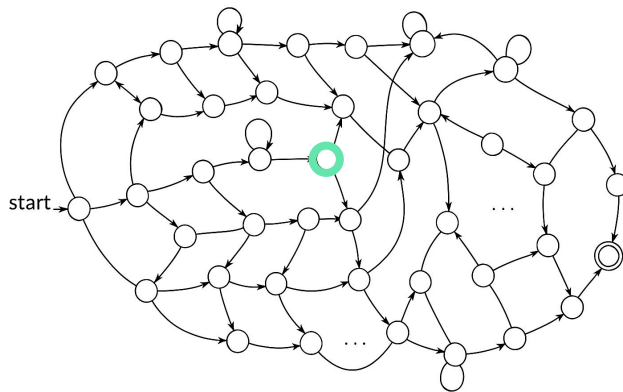
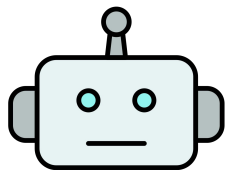
From single-task to multi-task

In a multi-task setting, we do not know the automaton beforehand



From single-task to multi-task

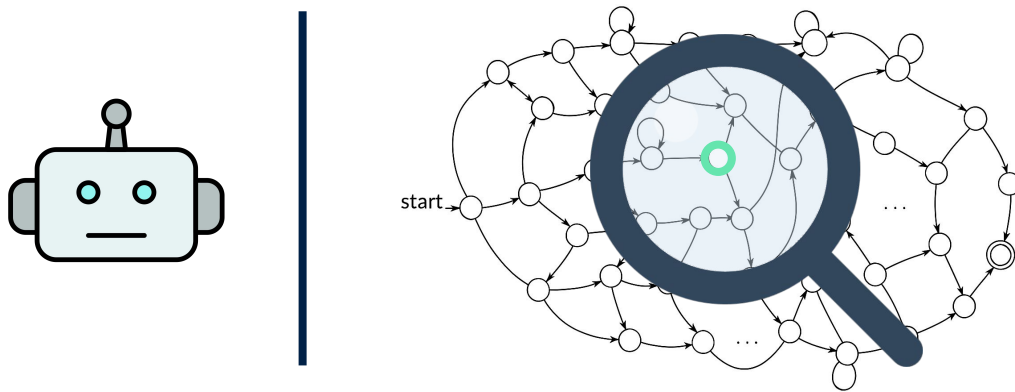
In a multi-task setting, we do not know the automaton beforehand



What is a **general representation** of the automaton state that can be used to condition the policy?

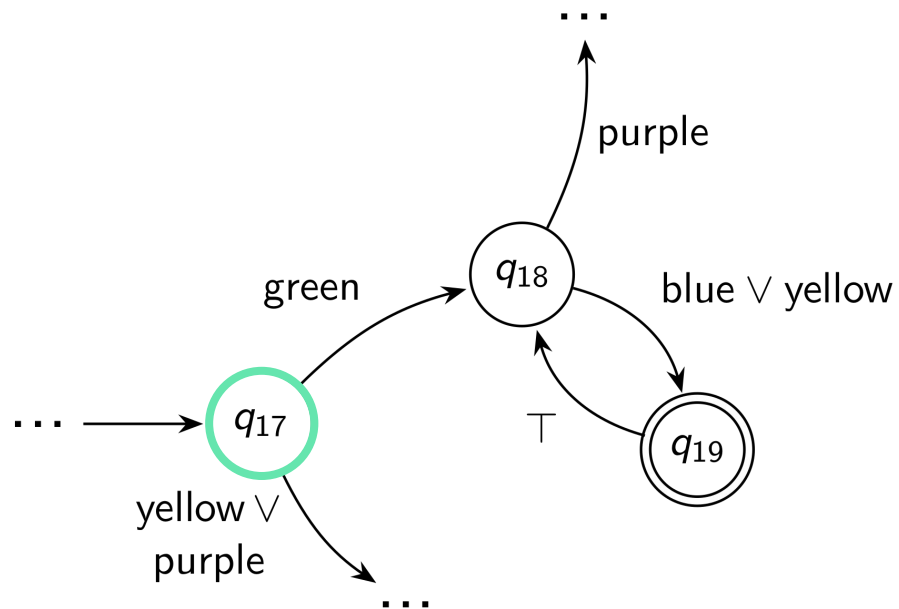
From single-task to multi-task

In a multi-task setting, we do not know the automaton beforehand

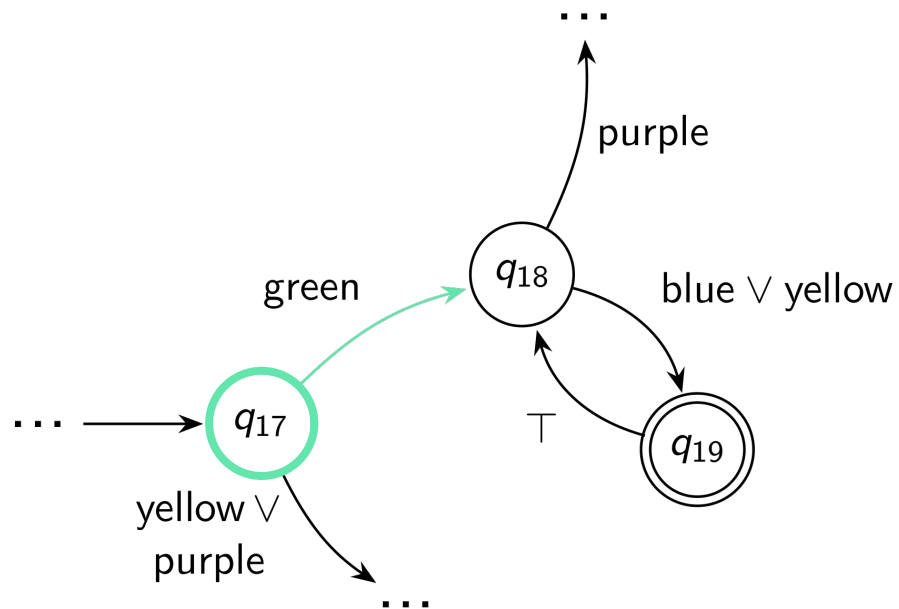


What is a **general representation** of the automaton state that can be used to condition the policy?

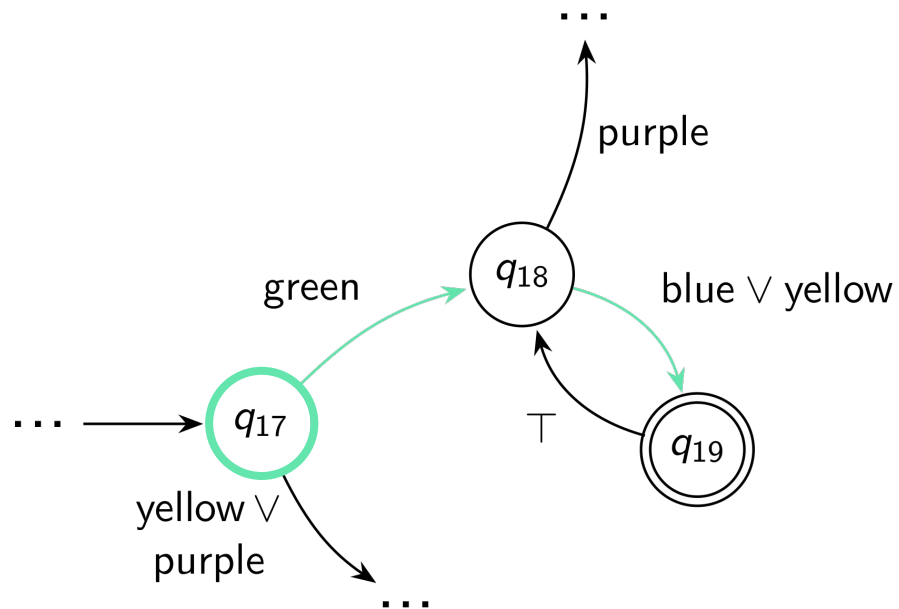
Reach-avoid sequences



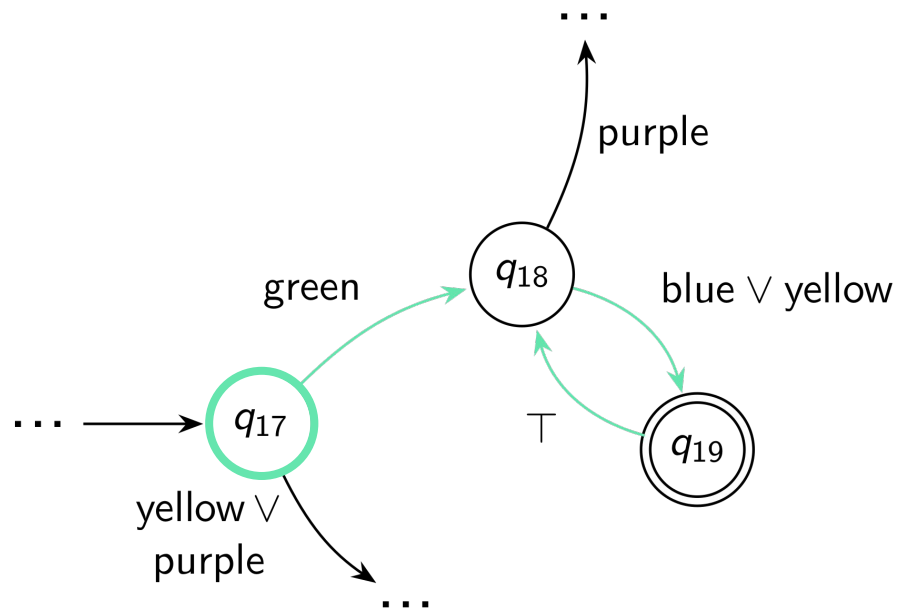
Reach-avoid sequences



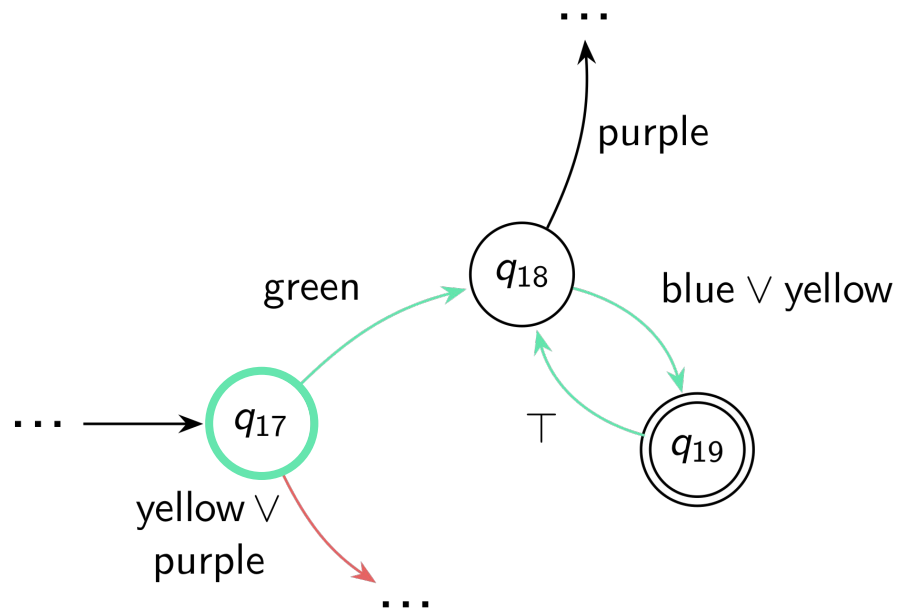
Reach-avoid sequences



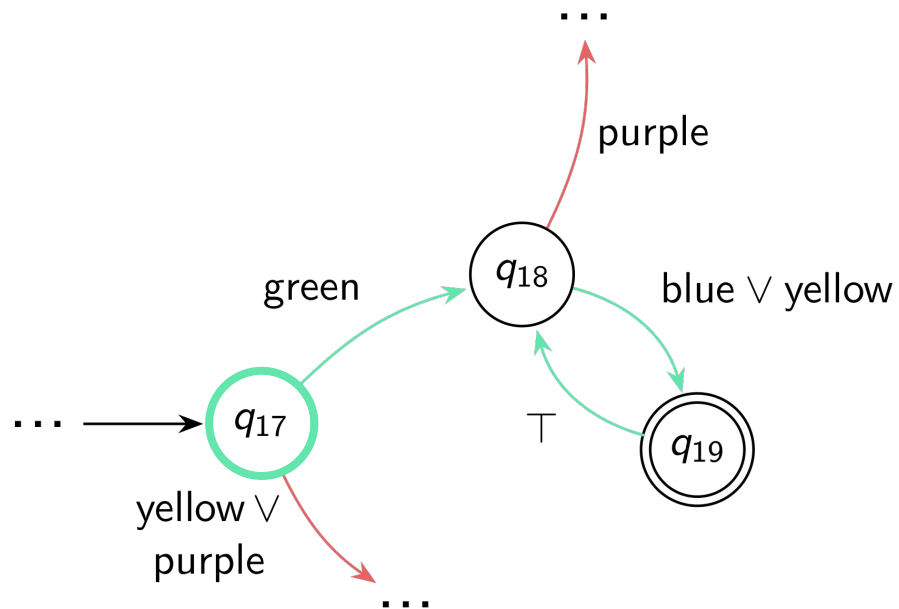
Reach-avoid sequences



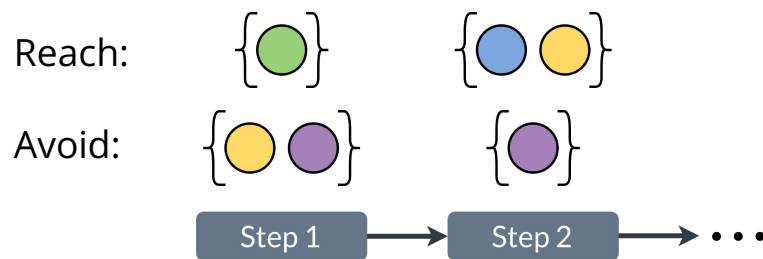
Reach-avoid sequences



Reach-avoid sequences

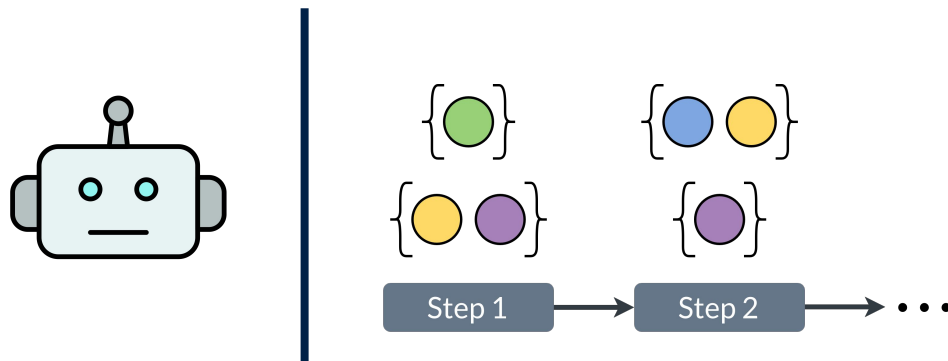


Reach-avoid sequences



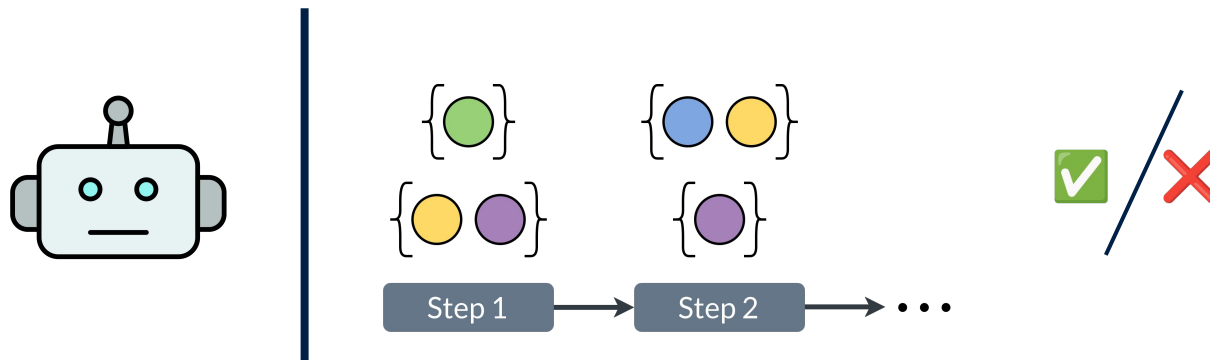
Training a general policy

We use goal-conditioned RL to train a general policy:



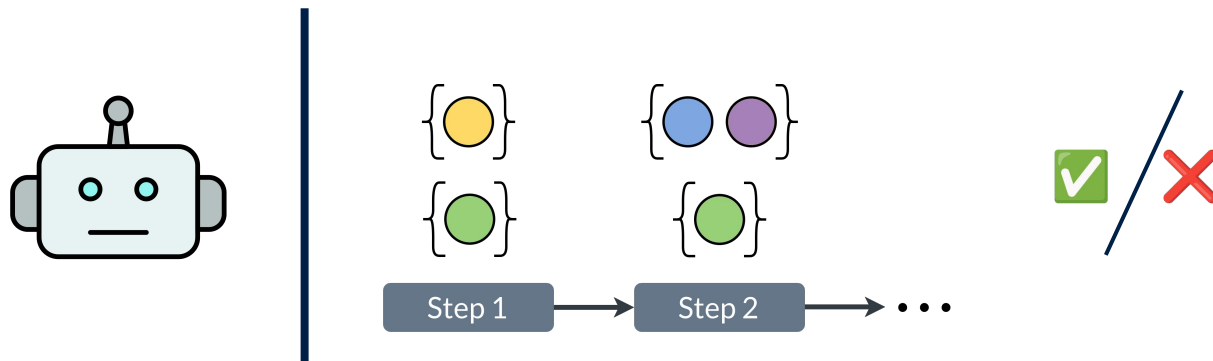
Training a general policy

We use goal-conditioned RL to train a general policy:



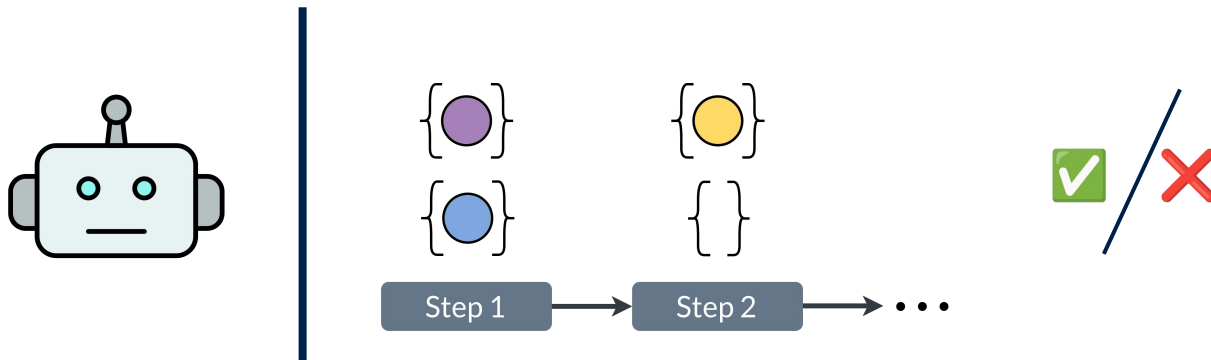
Training a general policy

We use goal-conditioned RL to train a general policy:



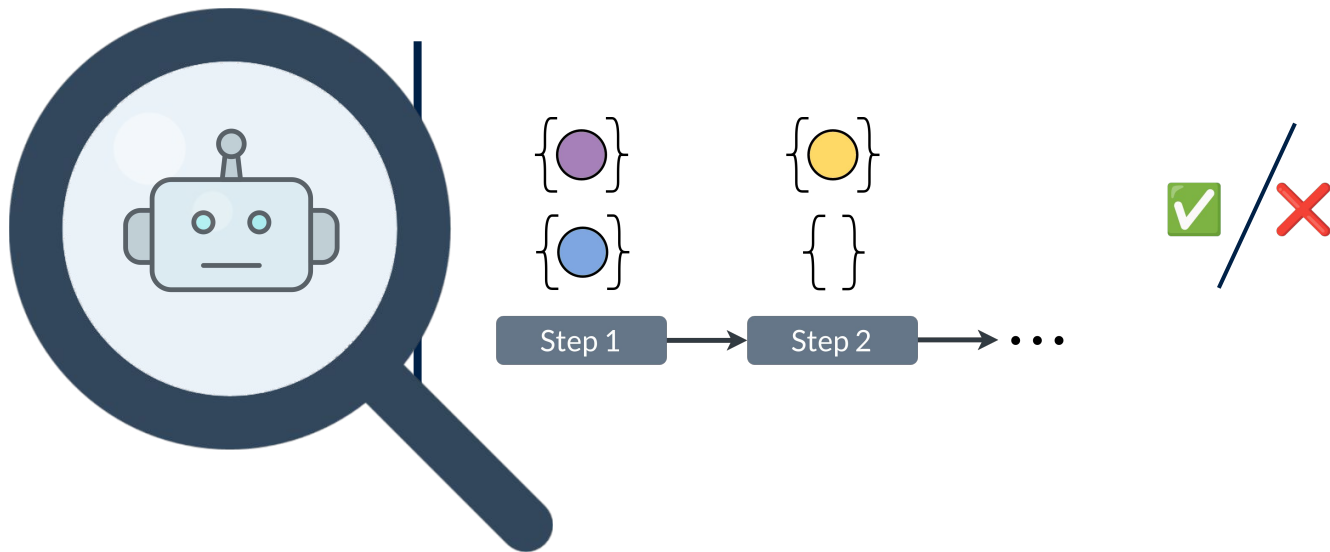
Training a general policy

We use goal-conditioned RL to train a general policy:

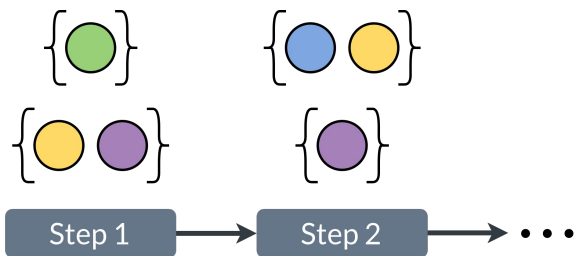


Training a general policy

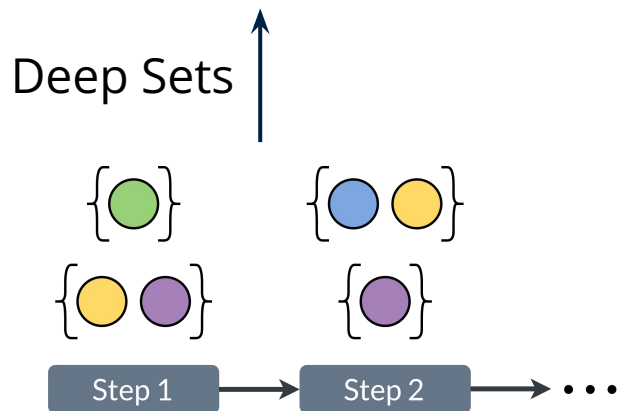
We use goal-conditioned RL to train a general policy:



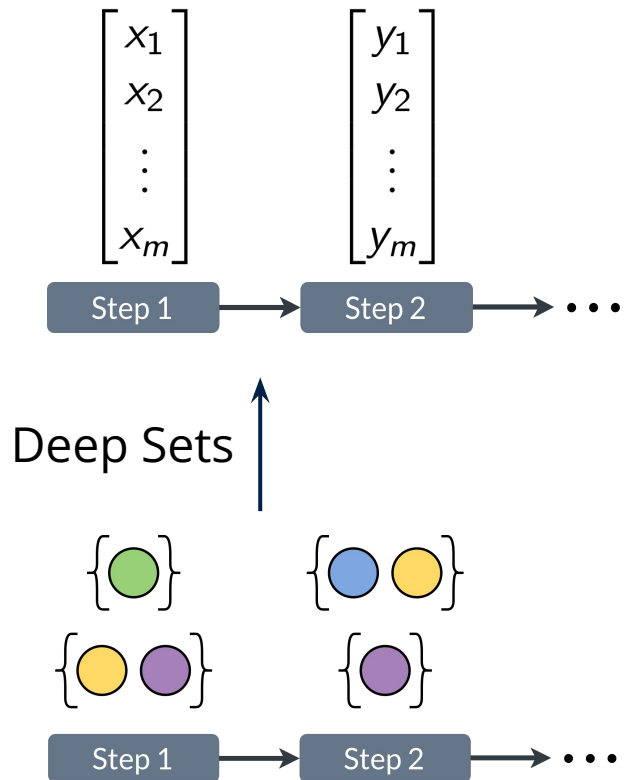
Model architecture



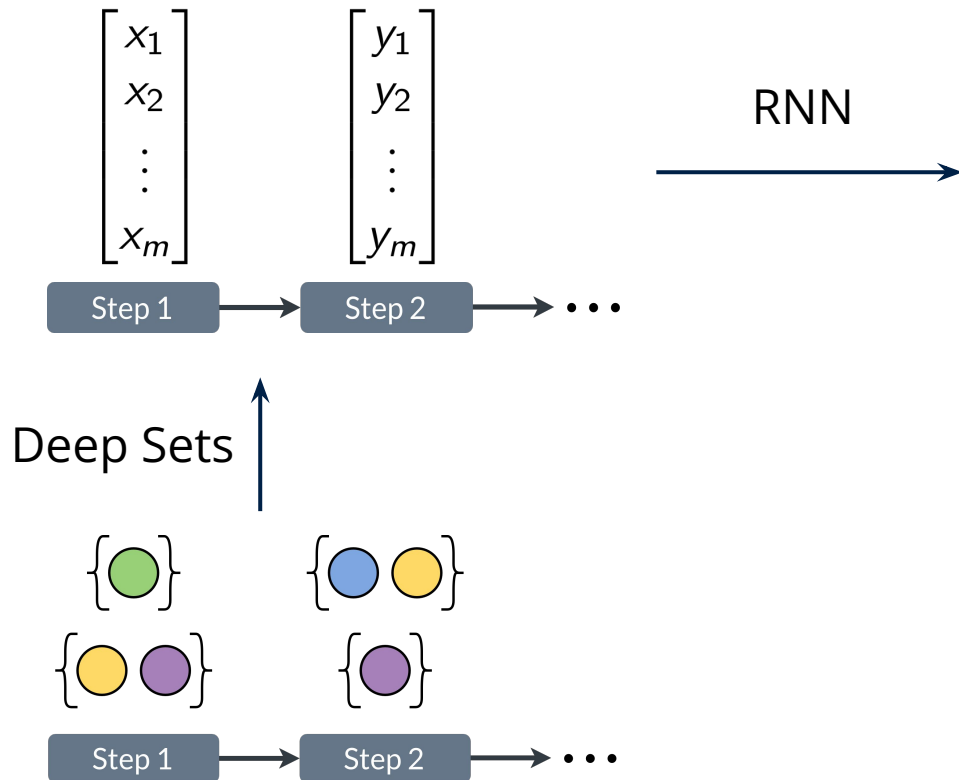
Model architecture



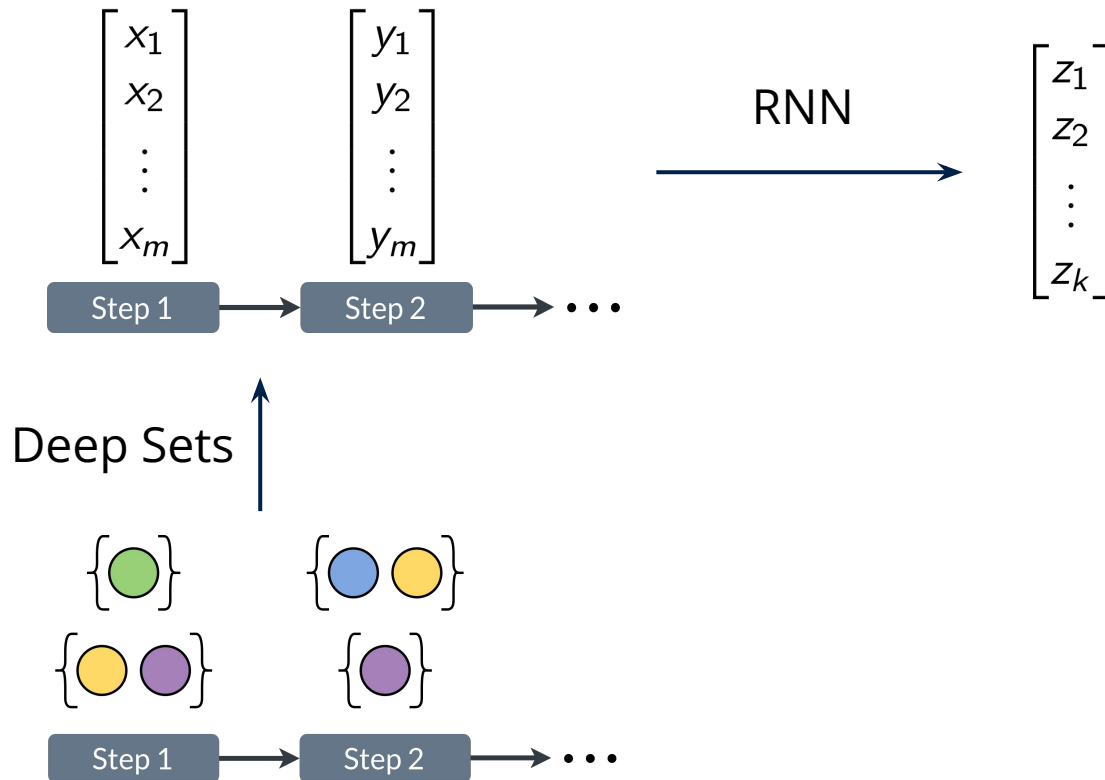
Model architecture



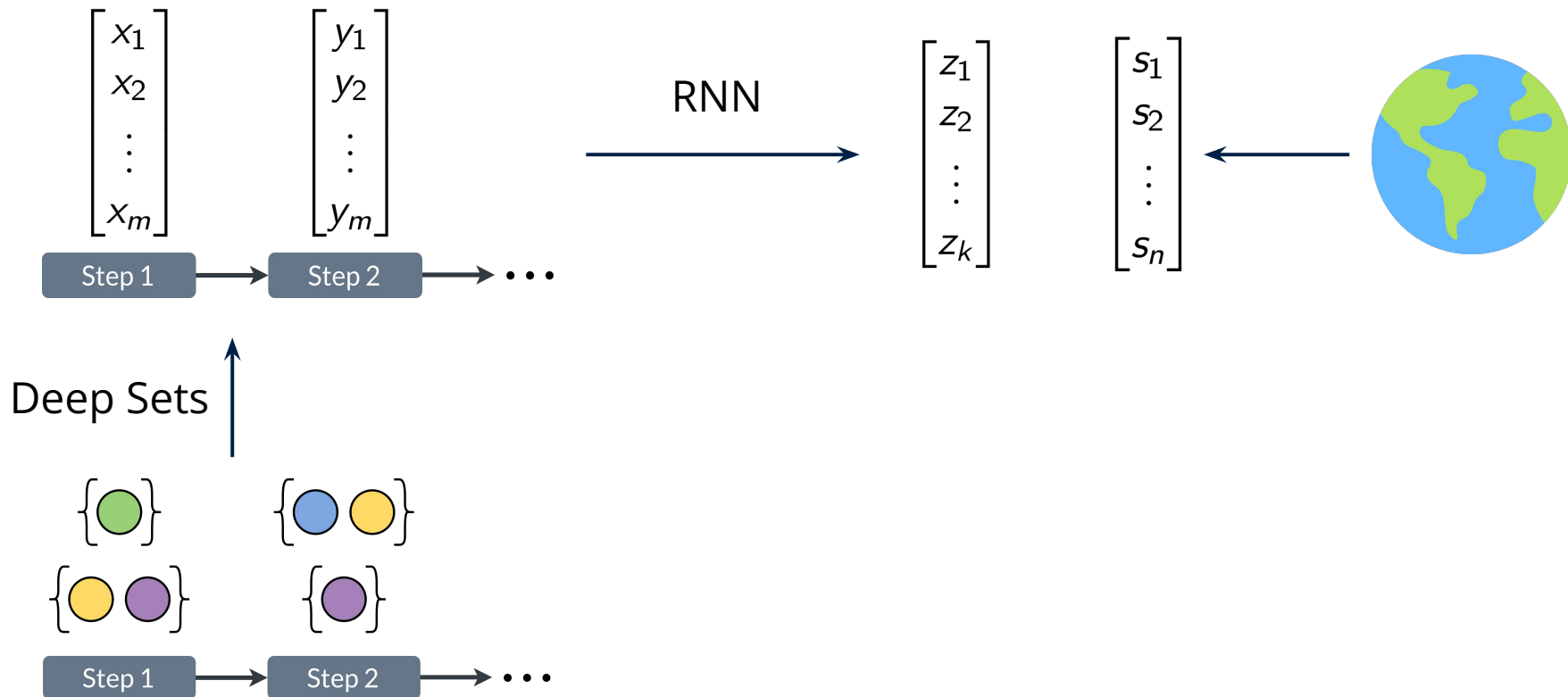
Model architecture



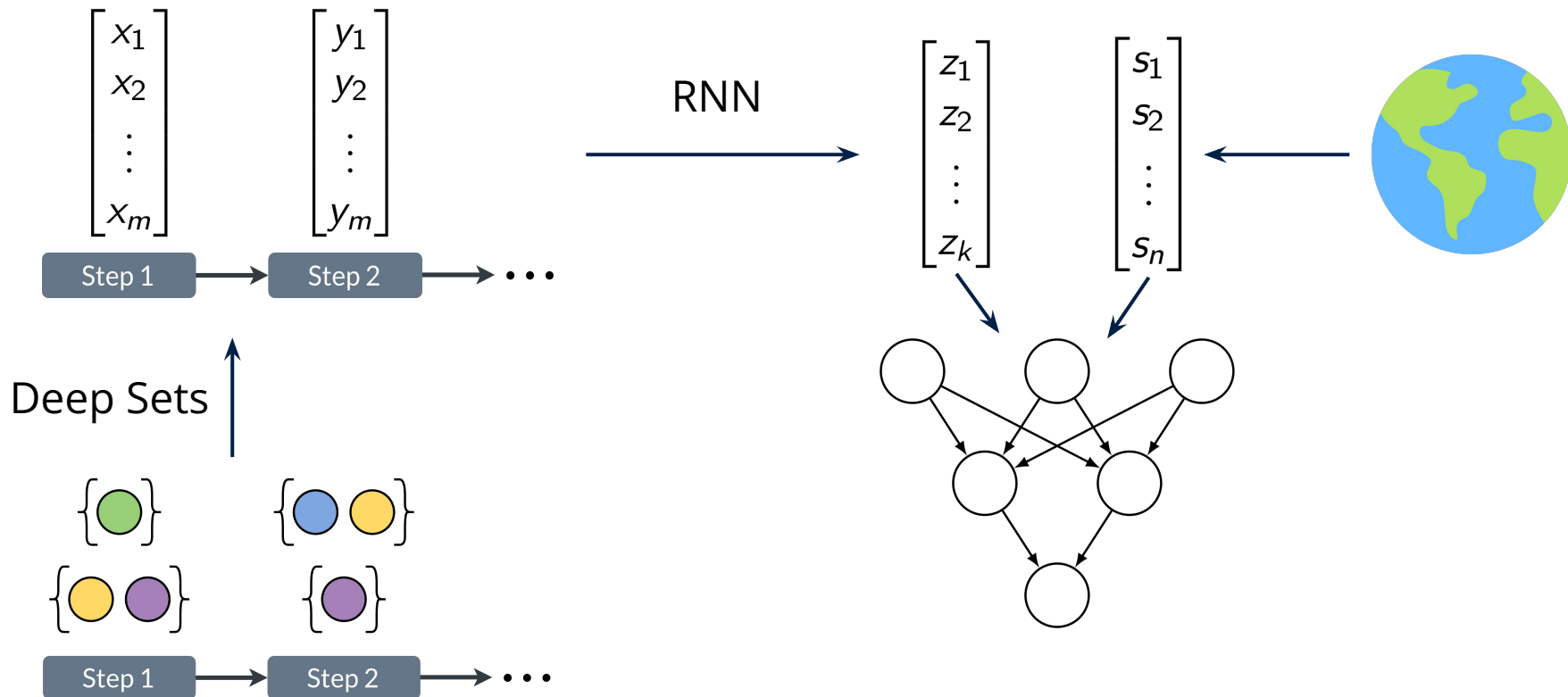
Model architecture



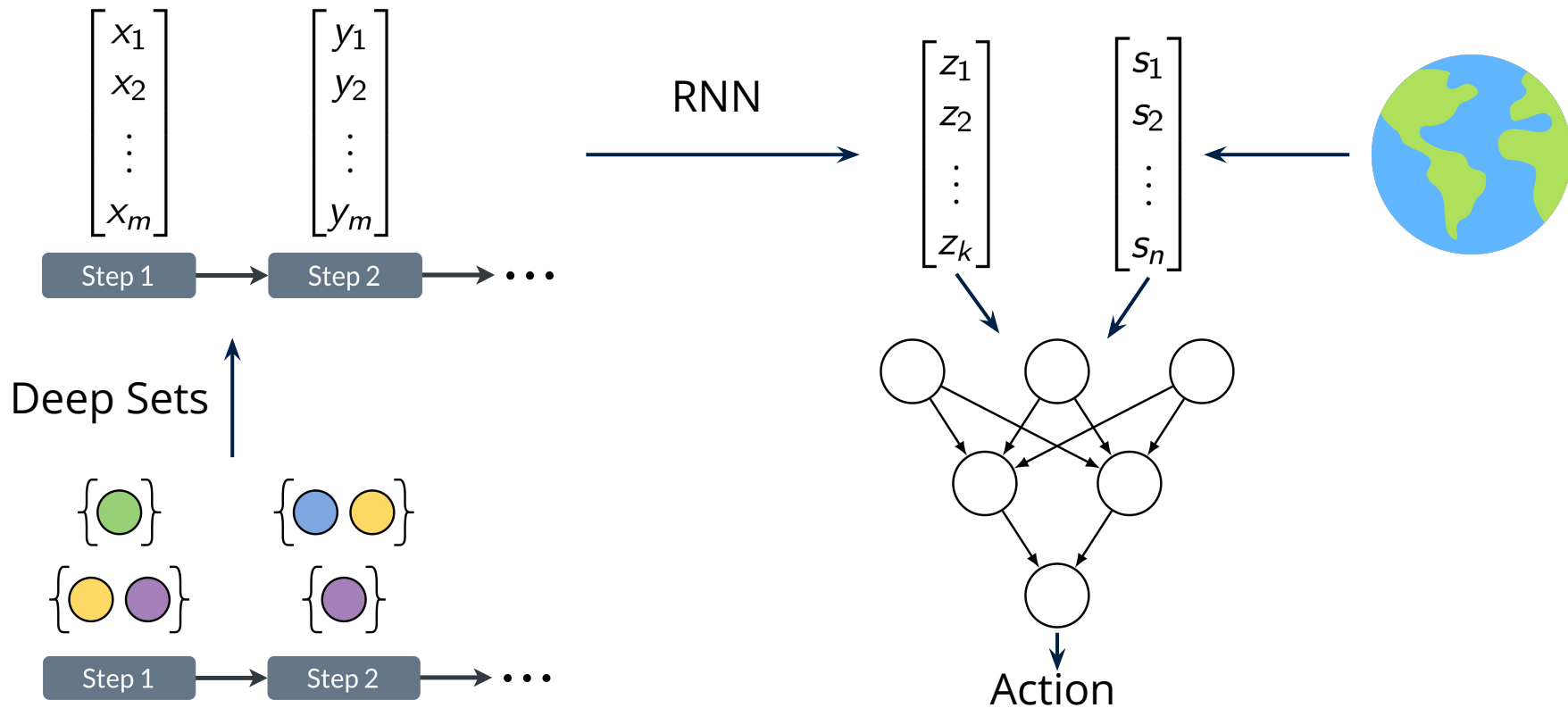
Model architecture



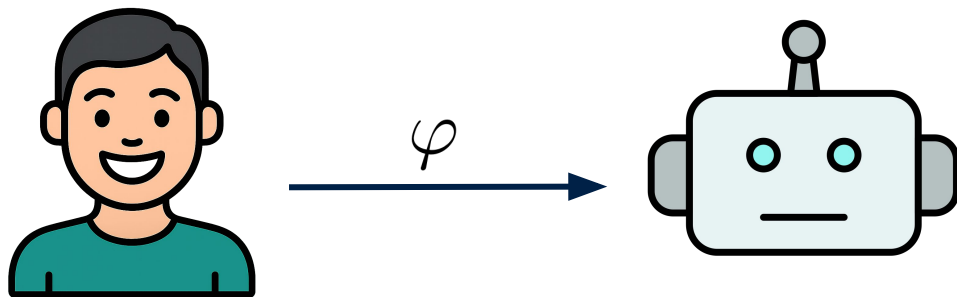
Model architecture



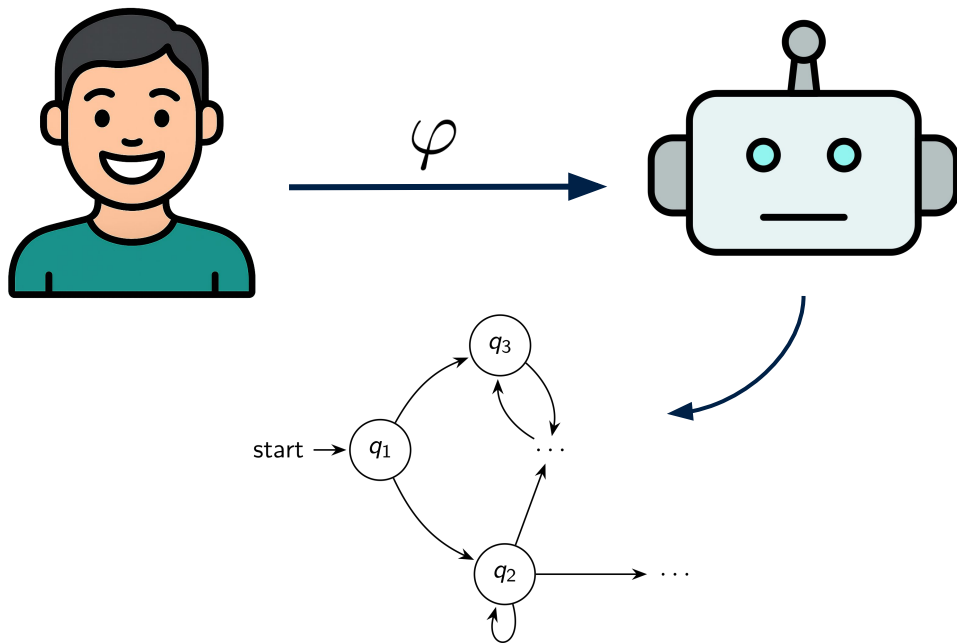
Model architecture



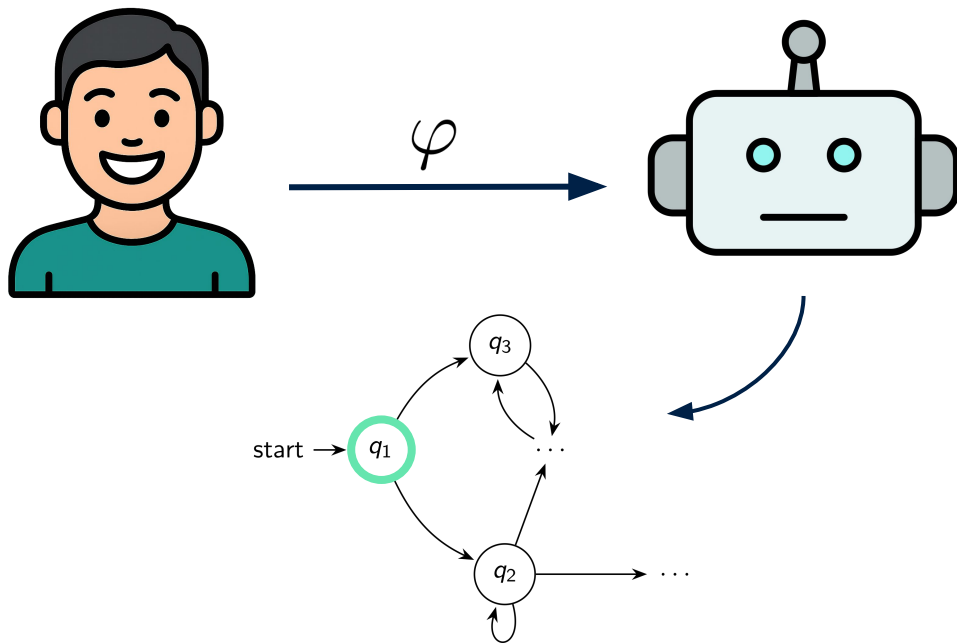
Test-time policy execution



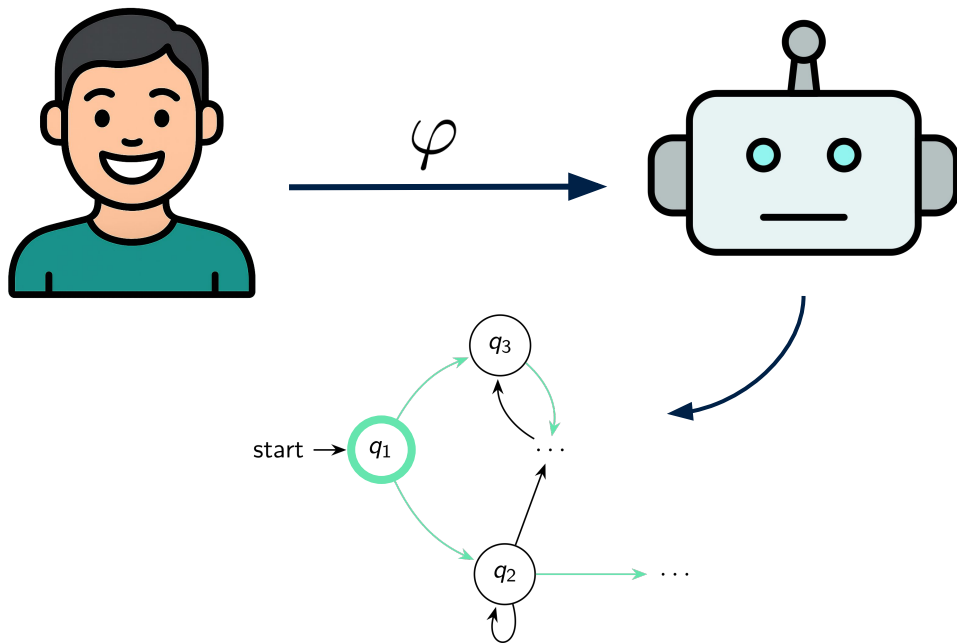
Test-time policy execution



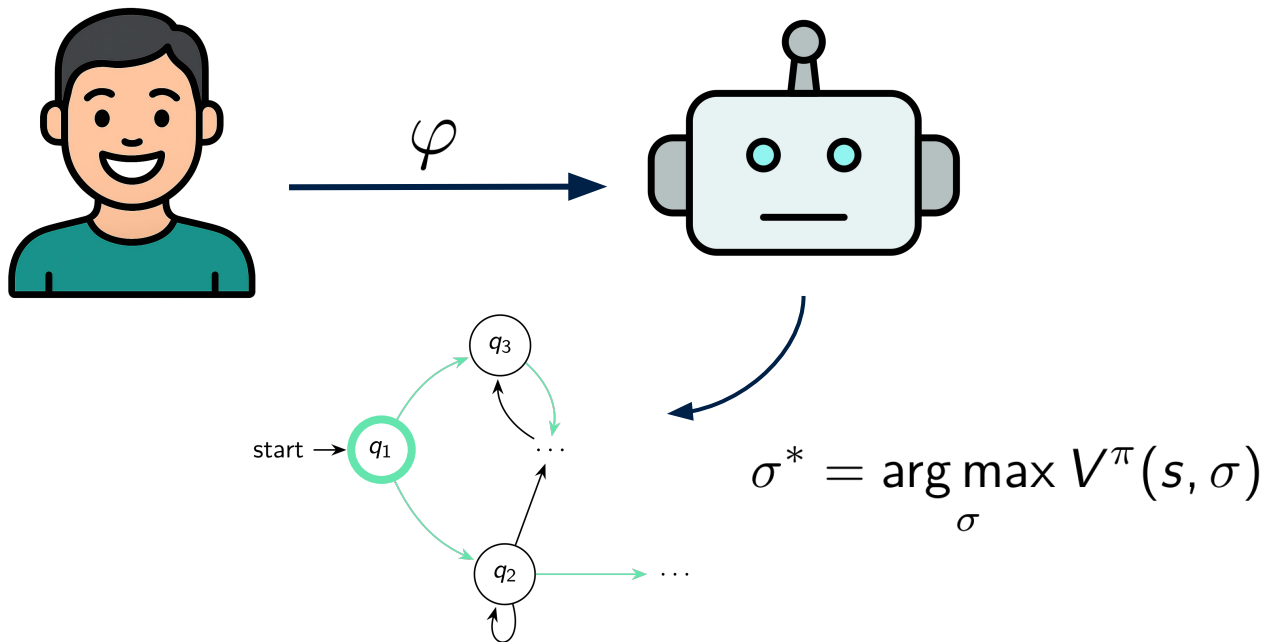
Test-time policy execution



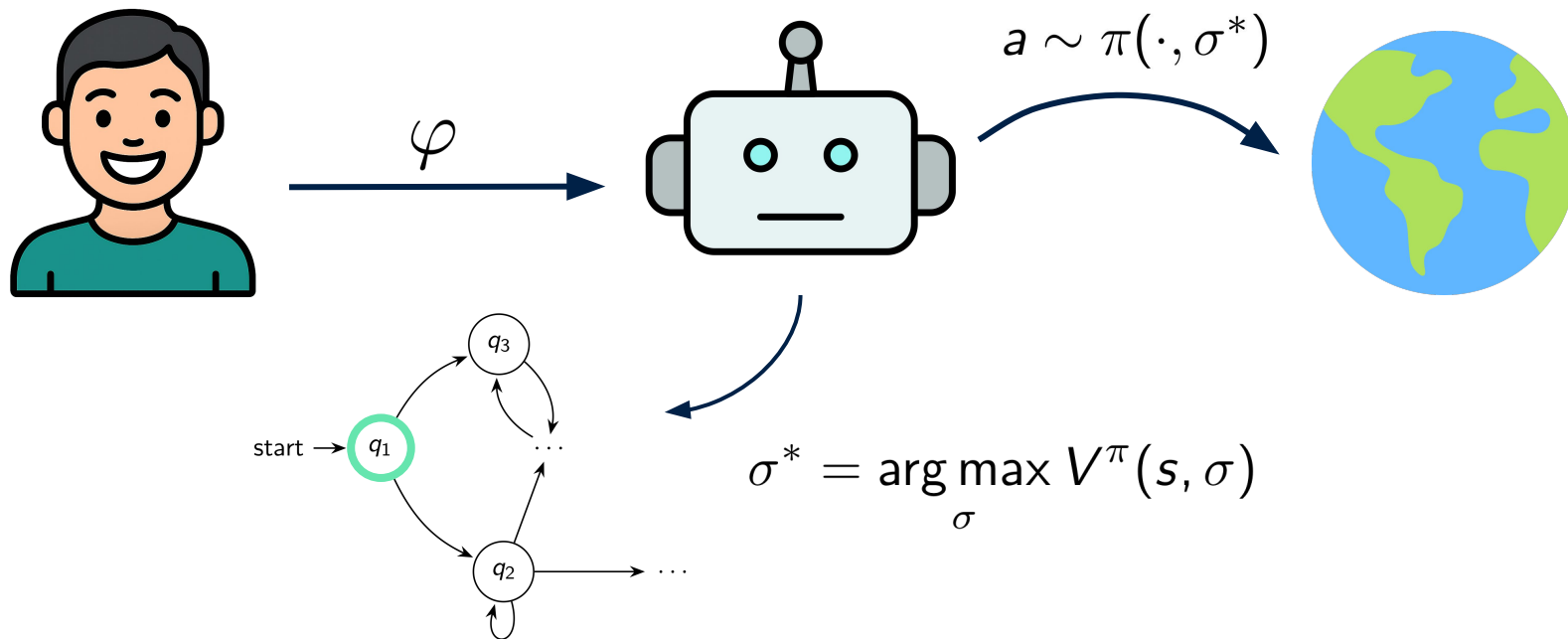
Test-time policy execution



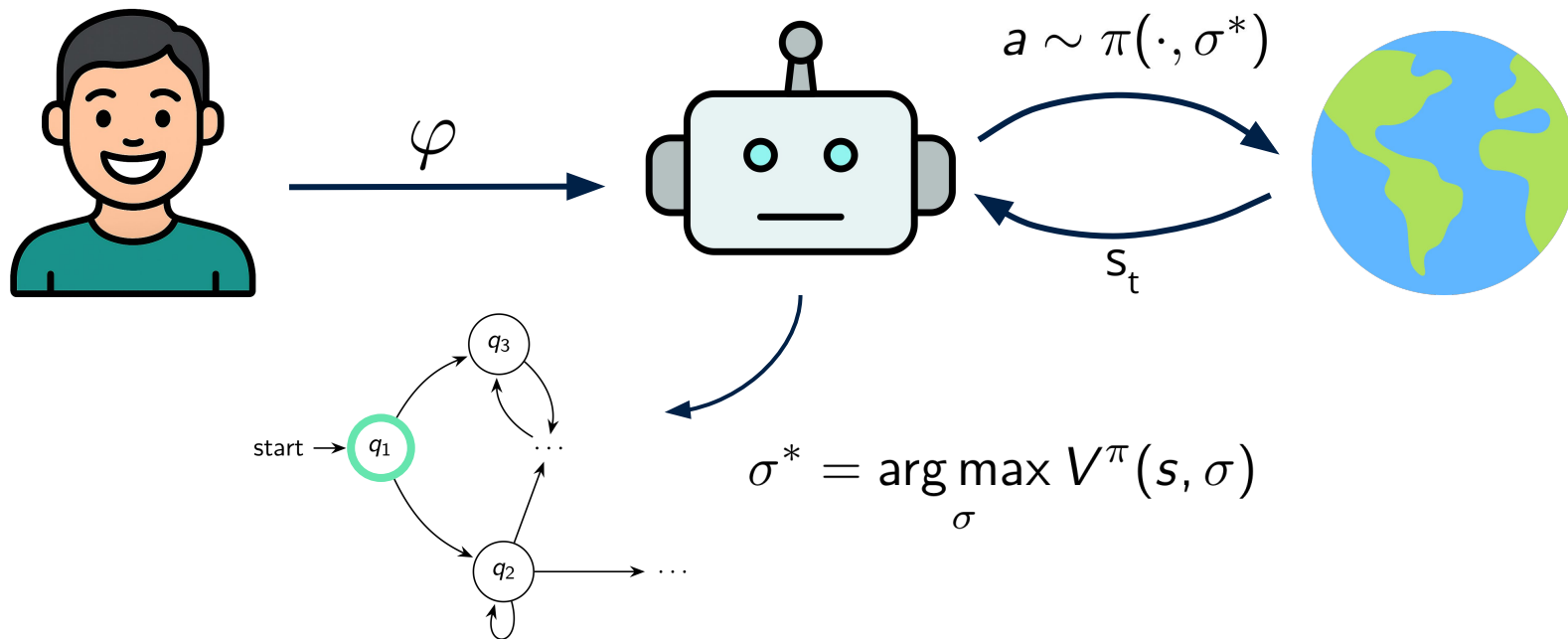
Test-time policy execution



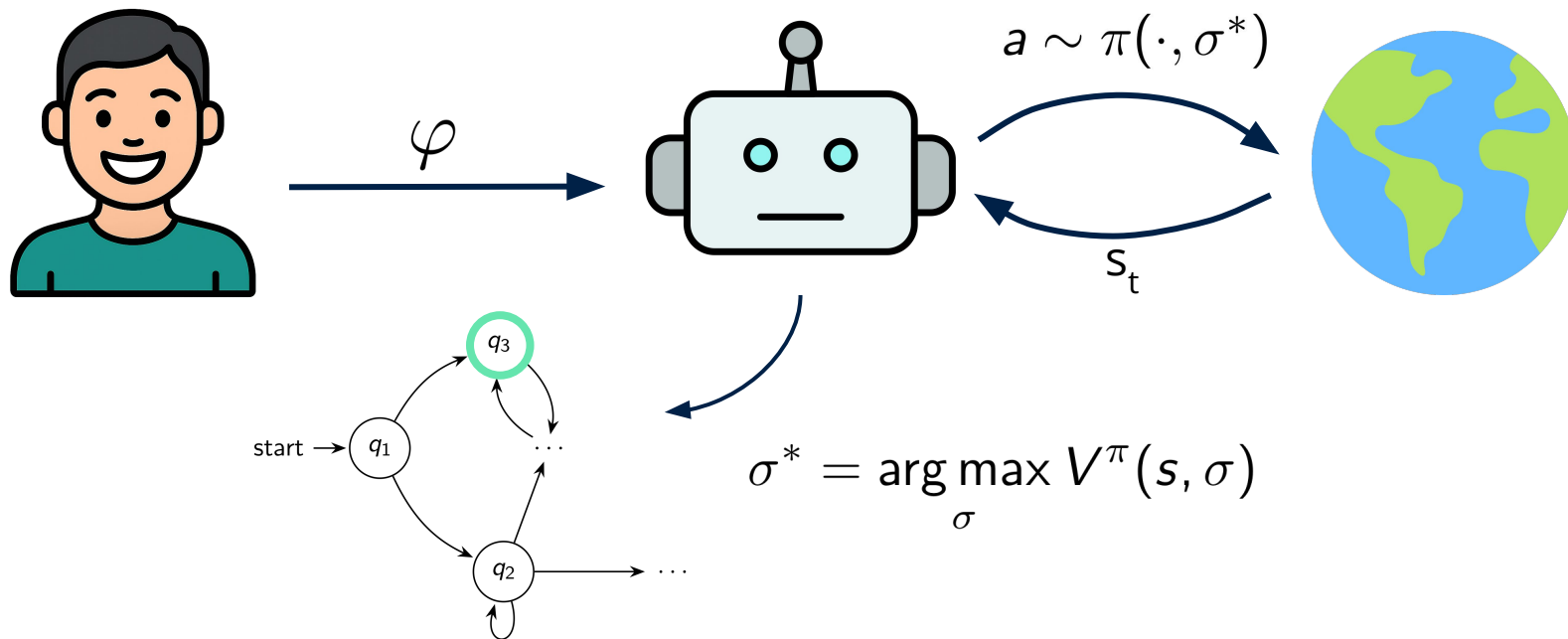
Test-time policy execution



Test-time policy execution

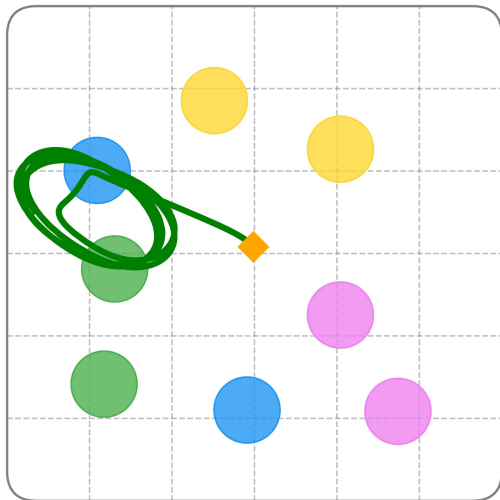


Test-time policy execution

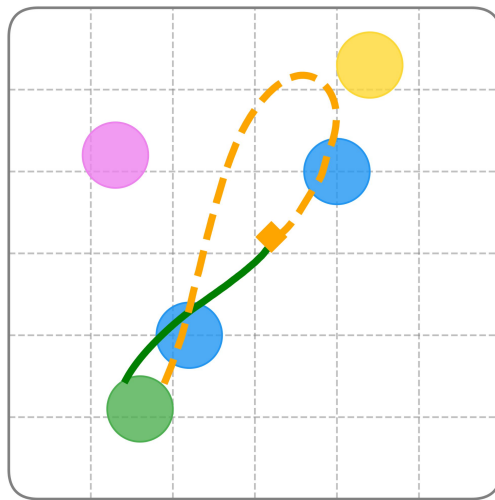


Discussion & Results

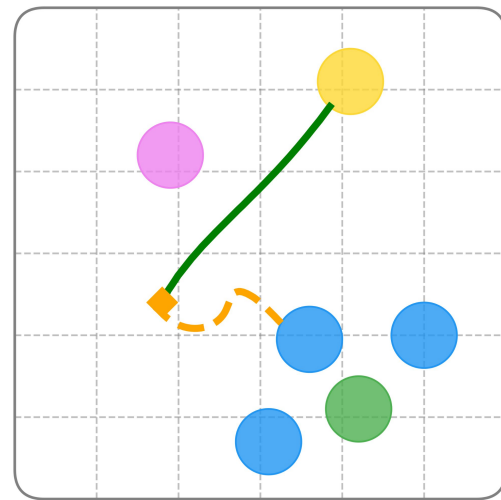
Discussion



Infinite-horizon
tasks

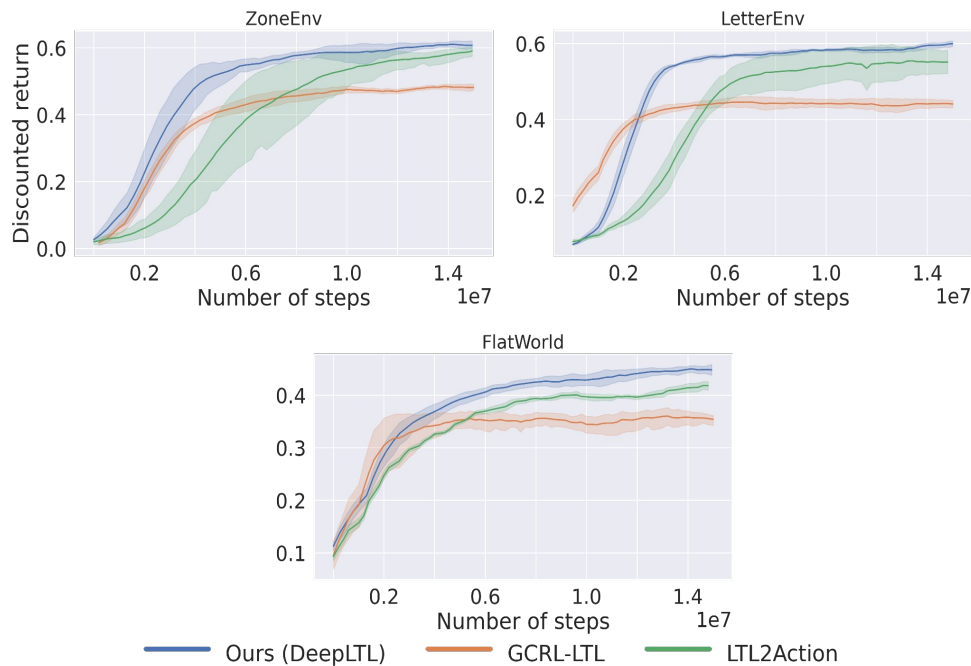


Optimality



Safety

Results



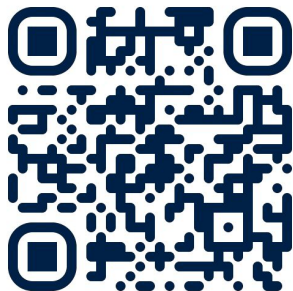
		LTL2Action	GCRL-LTL	DeepLTL
LetterWorld	φ_1	0.75 ± 0.18	0.94 ± 0.05	1.00 ± 0.00
	φ_2	0.79 ± 0.10	0.94 ± 0.03	0.98 ± 0.00
	φ_3	0.41 ± 0.14	1.00 ± 0.00	1.00 ± 0.00
	φ_4	0.72 ± 0.17	0.82 ± 0.07	0.97 ± 0.01
	φ_5	0.44 ± 0.26	1.00 ± 0.00	1.00 ± 0.00
ZoneEnv	φ_6	0.60 ± 0.20	0.85 ± 0.03	0.92 ± 0.06
	φ_7	0.14 ± 0.18	0.85 ± 0.05	0.91 ± 0.03
	φ_8	0.67 ± 0.26	0.89 ± 0.04	0.96 ± 0.04
	φ_9	0.69 ± 0.22	0.87 ± 0.02	0.90 ± 0.03
	φ_{10}	0.66 ± 0.19	0.85 ± 0.02	0.91 ± 0.02
	φ_{11}	0.93 ± 0.07	0.89 ± 0.01	0.98 ± 0.01
FlatWorld	φ_{12}	1.00 ± 0.00	0.82 ± 0.41	1.00 ± 0.00
	φ_{13}	0.63 ± 0.50	0.00 ± 0.00	1.00 ± 0.00
	φ_{14}	0.71 ± 0.40	0.73 ± 0.41	0.98 ± 0.01
	φ_{15}	0.07 ± 0.02	0.73 ± 0.03	0.86 ± 0.01
	φ_{16}	0.56 ± 0.35	0.64 ± 0.08	1.00 ± 0.01

Further resources

Website: deep-ltl.github.io

arXiv: arxiv.org/abs/2410.04631

GitHub: [mathiasj33/deep-ltl](https://github.com/mathiasj33/deep-ltl)



 mathias-jackermeier.me

 mathias.jackermeier@cs.ox.ac.uk

 [mathiasj33](https://github.com/mathiasj33)

 [@m_jackermeier](https://twitter.com/m_jackermeier)