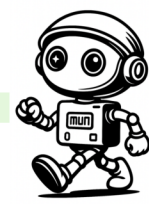
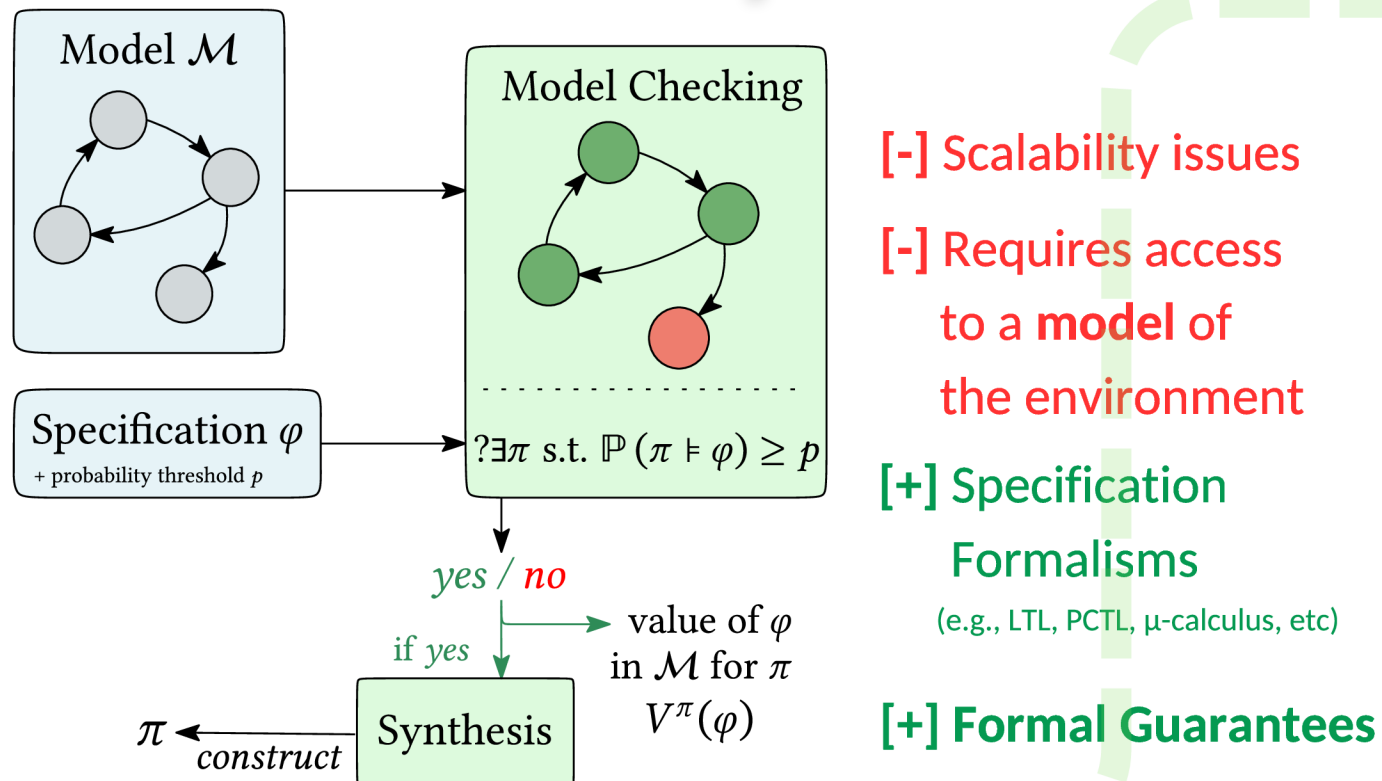


# Composing Reinforcement Learning Policies, with Formal Guarantees



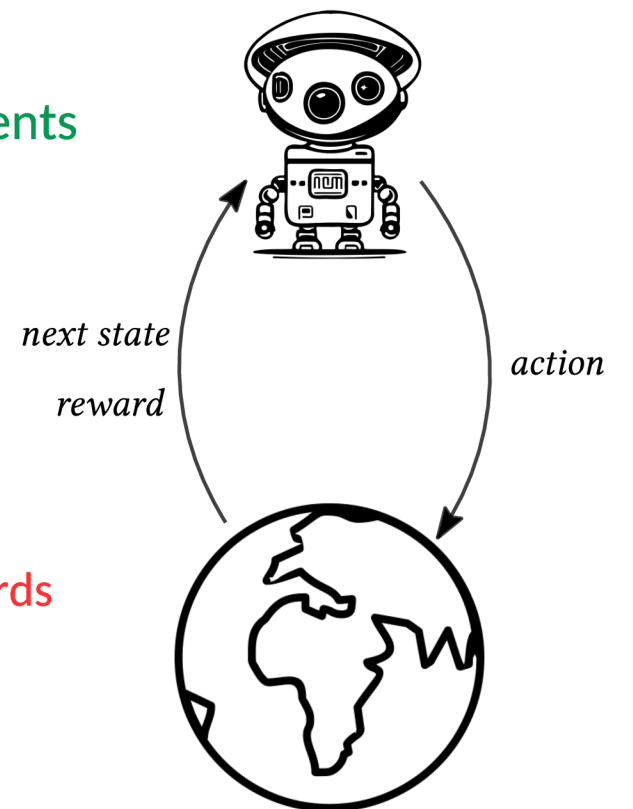
Florent Delgrange, Guy Avni, Christian Schilling, Anna Lukina, Ann Nowé, Guillermo A. Pérez

## Reactive Synthesis

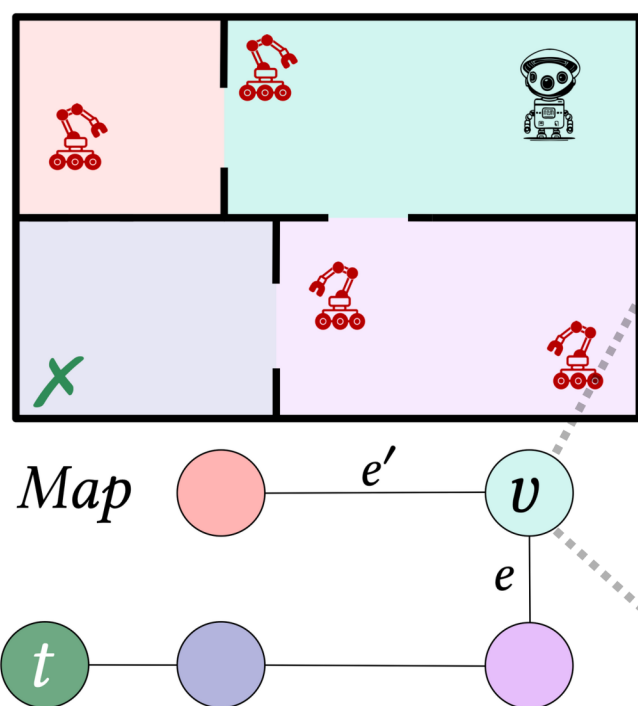


## Reinforcement Learning

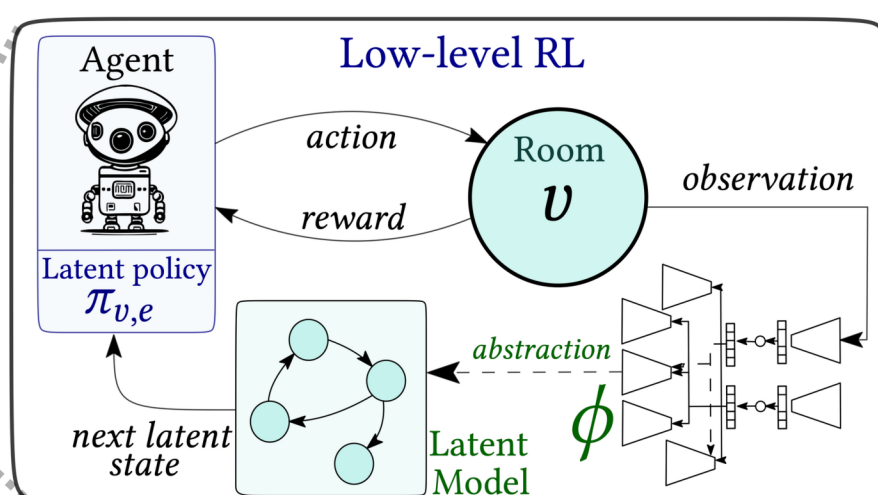
- [+] Scales to complex environments**
- [+] Interacting with the environment is sufficient**
- [-] Requires designing a reward function**  
→ reward engineering, sparse rewards
- [-] No Guarantee**



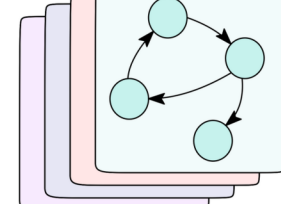
### Environment



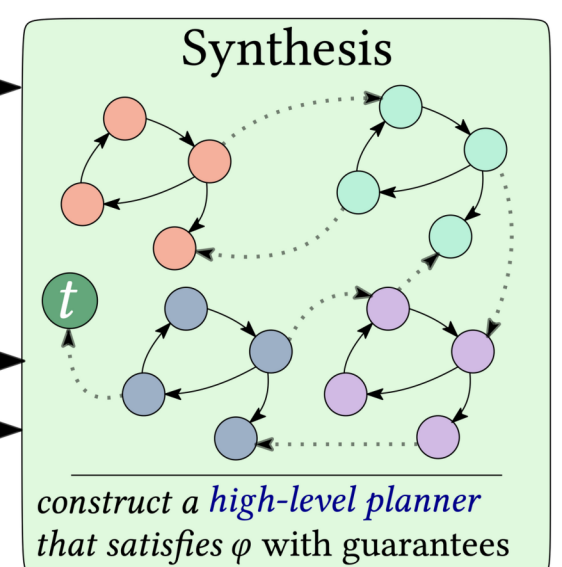
Our goal: get the best **[+]** of both worlds



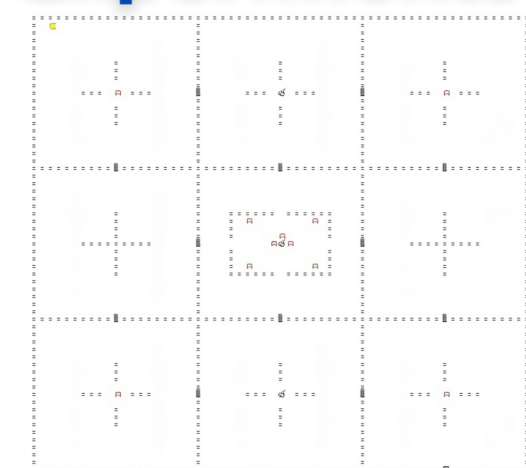
low-level latent models/policies (one per room, direction)



Specification  $\varphi$ : "reach  $t$ "



## Experiments

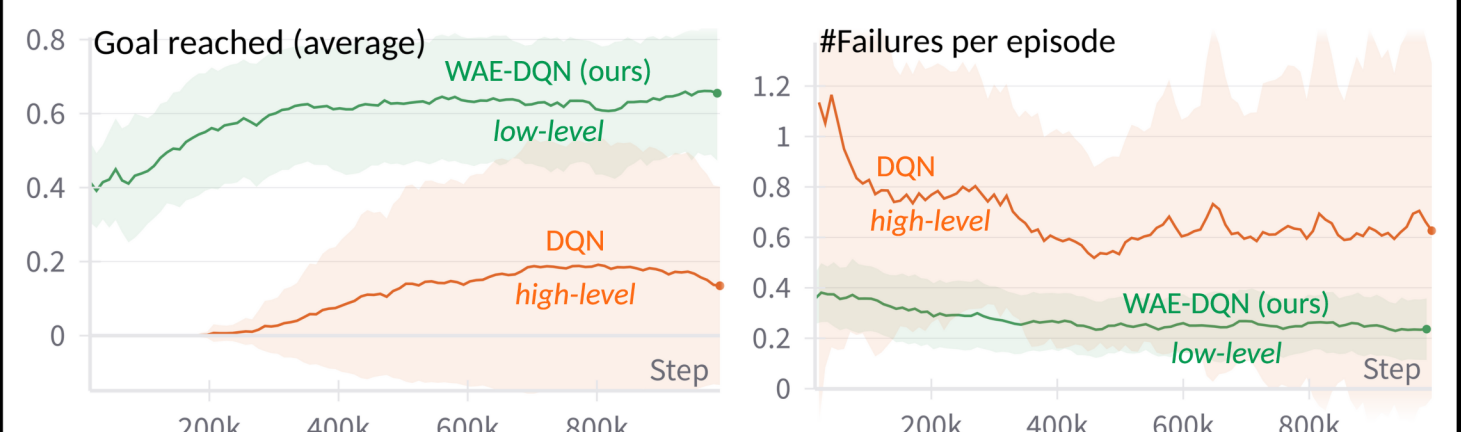


N: #rooms, LP: life points, A: #adversaries

|            | N  | LP | A  | avg. return ( $\gamma = 1$ ) | latent value | avg. value (original)   |
|------------|----|----|----|------------------------------|--------------|-------------------------|
| Grid World | 9  | 1  | 11 | $0.5467 \pm 0.1017$          | 0.1378       | $0.07506 \pm 0.01664$   |
|            | 9  | 3  | 11 | $0.7 \pm 0.09428$            | 0.4343       | $0.01 \pm 0.00163$      |
|            | 25 | 3  | 23 | $0.4933 \pm 0.09832$         | 0.1763       | $0.007833 \pm 0.002131$ |
|            | 25 | 5  | 23 | $0.5667 \pm 0.07817$         | 0.346        | $0.00832 \pm 0.00288$   |
| ViZDoom    | 49 | 7  | 47 | $0.02667 \pm 0.01491$        | 0.004229     | $5.565e-6 \pm 7e-6$     |
|            | 8  | /  | 8  | $0.89333 \pm 0.059628$       | 0.24171      | $0.23405 \pm 0.014781$  |
|            | 8  | /  | 14 | $0.78 \pm 0.064979$          | 0.16459      | $0.16733 \pm 0.023117$  |
|            | 8  | /  | 20 | $0.39333 \pm 0.11643$        | 0.086714     | $0.06898 \pm 0.017788$  |

|   | Grid World | ViZDoom |
|---|------------|---------|
| → | 0.50412    | 0.32011 |
| ← | 0.77787    | 0.44883 |
| ↑ | 0.49631    | 0.37931 |
| ↓ | 0.48058    | 0.48108 |

Table 2: PAC bounds.



## Formal Guarantees

- Learn  $\bar{\mathcal{M}}$  via a loss:

$$L_P^{v,e} = \mathbb{E}_{s,a \sim \pi_{v,e}} \text{TV}(\phi_{\#} \mathbf{P}(\cdot | s, a), \bar{\mathbf{P}}(\cdot | \phi(s), a))$$

- Room-wise **Abstraction Quality Guarantees**:

$$\left| V_v^{\pi_{v,e}}(\varphi) - V_{\bar{\mathcal{M}}}^{\pi_{v,e}}(\varphi) \right| \leq \text{AEL}(\pi_{v,e}) \cdot \frac{L_P^{v,e}}{1-\gamma}$$

Average Episode Length

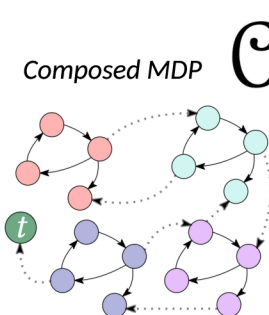
PAC Learnable

- Memory is **necessary** for the HL planner  $\tau$

- Optimal  $\tau$  obtained by solving an MDP

- Lifting the guarantees** to the high-level:

$$\left| V_{\text{true}}^{\tau}(\varphi) - V_{\bar{\mathcal{C}}}^{\tau}(\varphi) \right| \leq \text{AEL}(\tau) \cdot \frac{L_I + K \cdot \mathbb{E}_{(v,e) \sim \tau} [L_P^{v,e}]}{1-\gamma}$$



check out the paper!

Paper



Blogpost



in/florent-delgrange



@florentdelgrange.bsky.social



@f\_delgrange



<https://delgrange.me/>

