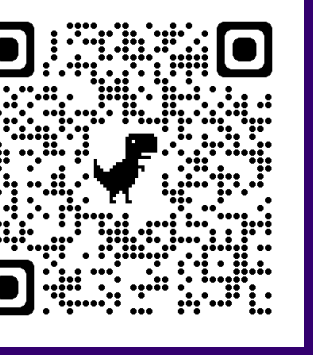




Preliminaries

Markov Decision Processes

- S states
- A actions
- Planning horizon H
- Reward function $R(s, a)$, such that for any possible trajectory, its total reward in a single episode is upper-bounded by 1
- Transition probability $P(s'|s, a)$
- Initial state distribution $\nu(s)$
- Maximum transition support $\Gamma = \max_{s,a} \|P(\cdot|s, a)\|_0$

Latent MDPs

- A distribution of M MDPs $\mathcal{M} = \{\mathcal{M}_1, \dots, \mathcal{M}_M\}$
- Each with weight (probability) $w_1, \dots, w_M$
- All MDPs share the same states, actions and horizon
- But have their own reward function R_m , transition P_m and initial state distribution ν_m

Policy and alpha vectors

- $\pi : \mathcal{H} \rightarrow \mathcal{A}$, where $\mathcal{H}$ is the set of all histories
- Alpha vectors (generalization of value and Q -functions)

$$\alpha_m^\pi(h) := \mathbb{E}_{\pi, \mathcal{M}_m} \left[\sum_{t'=t}^H R_m(s_{t'}, a_{t'}) \mid h_t = h \right],$$

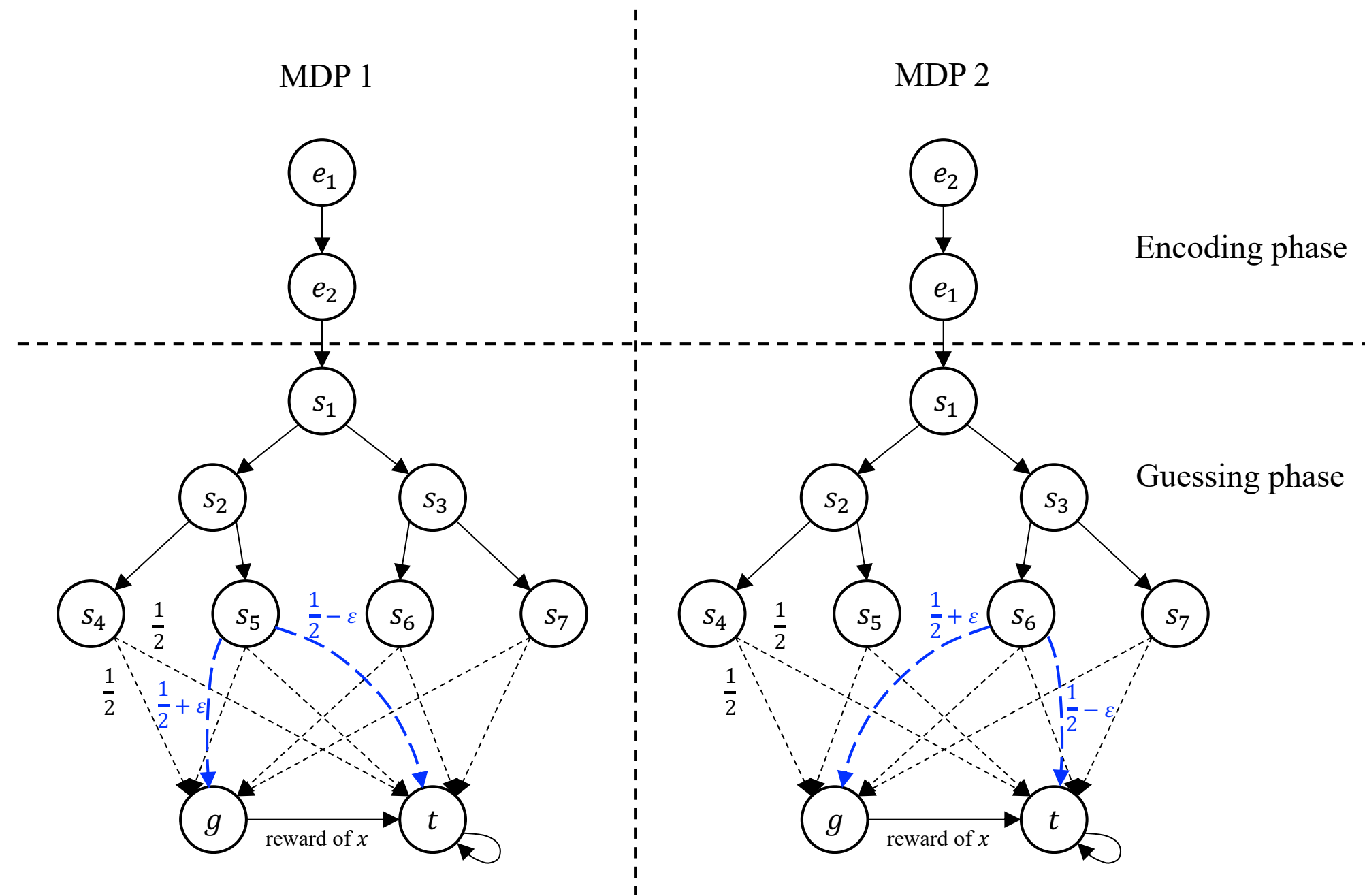
$$\alpha_m^\pi(h, a) := \mathbb{E}_{\pi, \mathcal{M}_m} \left[\sum_{t'=t}^H R_m(s_{t'}, a_{t'}) \mid (h_t, a_t) = (h, a) \right].$$

- Value function $V^\pi = \sum_{m,s} w_m \nu_m(s) \alpha_m^\pi(s)$

Episodic reinforcement learning with context in hindsight

- K episodes, with policies $\pi^1, \pi^2, \dots, \pi^K$
- **Context in hindsight:** At the beginning, sample $m \sim \{w\}$, but only tell the agent m when the episode ends
- Performance measure

$$\text{Regret}(K) := \sum_{k=1}^K (V^* - V^{\pi^k}).$$



Motivation

A problem harder than MDPs while easier than POMDPs

- LMDPs are collections of MDPs with a hidden context
- LMDPs are POMDPs with invariant unobservables throughout an episode: $s = (m, o) \rightarrow s' = (m, o')$

RL environments have different randomness:

We desire a regret that:

- Better than a previous work of $\tilde{O}(\sqrt{MHS^2AK})$
- Reduce to better guarantee if the LMDP is special, e.g., a deterministic MDP

Maximum Policy-Value Variance

For deterministic policy π , define

$$\text{Var}^\pi := \mathbb{V}(w \circ \nu, \alpha^\pi(\cdot)) + \mathbb{E}_\pi \left[\sum_{t=1}^H \mathbb{V}(P_m(\cdot|s_t, a_t), \alpha_m^\pi(h_t a_t r_t)) \right].$$

and $\text{Var}^* := \max_{\pi \in \Pi} \text{Var}^\pi$.

Discussion on Variances

- Under **Condition 1**, we have $\text{Var}_\tau^\Sigma \leq \tilde{O}(1)$ with high probability if τ is sampled by any policy and $\text{Var}^* \leq 1$. Without **Condition 1**, $\text{Var}_\tau^\Sigma \leq \tilde{O}(H^2)$ and $\text{Var}^* \leq H^2$.
- For deterministic MDPs, $\text{Var}_\tau^\Sigma = \text{Var}^* = 0$.
- $\text{Var}^* = 0 \implies \text{Var}_\tau^\Sigma = 0$, while reverse is not true.

Regret Upper Bound

We propose an algorithmic framework for solving LMDPs, which takes planning oracles as plug-in solvers. With two oracles we design (Bernstein confidence set and Monotonic Value Propagation for LMDP), we can guarantee a regret upper bound of

$$\text{Regret}(K) \leq \tilde{O}(\sqrt{\text{Var}^* M \Gamma S A K} + M S^2 A),$$

where $\tilde{O}$ hides polylog factors.

Remark:

- $\text{Var}^* \leq 1$ and $\Gamma \leq S$, so the worst case regret is $\tilde{O}(\sqrt{M S^2 A K} + M S^2 A)$, which is the first horizon-free bound for LMDPs
- If we view deterministic MDP as a special LMDP, we have $\text{Var}^* = 0$ and $M = 1$, regret is $\tilde{O}(S^2 A)$, which is a constant

Regret Lower Bound

For any variance level $0 < \mathcal{V} \leq O(1)$ and any algorithm π , there exists an LMDP $\mathcal{M}_\pi$ such that:

- $\text{Var}^* = \Theta(\mathcal{V})$;
- For $K \geq \tilde{\Omega}(M^2 + MSA)$, its expected regret in $\mathcal{M}_\pi$ after K episodes satisfies

$$\mathbb{E} \left[\sum_{k=1}^K (V^* - V^k) \mid \mathcal{M}_\pi, \pi \right] \geq \Omega(\sqrt{\mathcal{V} M S A K}).$$

High level idea: See the illustration in the top figure. We transform context in hindsight into context being told beforehand, while not affecting the optimal value function. Use a small portion of states to encode the context, then the optimal policy can extract information from them.