



Sharp Gap-Dependent Variance-Aware Regret Bounds for Tabular MDPs

Shulun Chen^{1*} Runlong Zhou² Zihan Zhang³ Maryam Fazal² Simon S. Du²

Preliminaries

Markov Decision Processes

- S states, A actions, planning horizon H
- Reward distribution $R_{s,a,h} \in \Delta([0, H])$
- Transition probability $P_{s,a,h}$

Policy and value functions

- $\pi = \{\pi_h\}_{h \in [H]}$ where $\pi_h : \mathcal{S} \rightarrow \mathcal{A}$, optimal π^*
- Value and Q functions

$$V_h^\pi(s) := \mathbb{E}^\pi \left[\sum_{t=h}^H r_t \mid s_h = s \right], V_0^\pi := \mathbb{E}^\pi \left[\sum_{t=1}^H r_t \right],$$

$$Q_h^\pi(s, a) := \mathbb{E}^\pi \left[\sum_{t=h}^H r_t \mid (s_h, a_h) = (s, a) \right].$$

- Optimal and Suboptimal Actions

$$\mathcal{Z}_{\text{opt}} = \{(s, a, h) \in \mathcal{S} \times \mathcal{A} \times [H] : \Delta_h(s, a) = 0\},$$

$$\mathcal{Z}_{\text{sub}} = \mathcal{S} \times \mathcal{A} \times [H] \setminus \mathcal{Z}_{\text{opt}}.$$

- Gap quantities

$$\Delta_h(s, a) := V_h^*(s) - Q_h^*(s, a), \Delta_{\min} := \min_{(s,a,h) \in \mathcal{Z}_{\text{sub}}} \Delta_h(s, a).$$

Episodic reinforcement learning

- K episodes, with policies $\pi^1, \pi^2, \dots, \pi^K$
- Performance measure

$$\text{Regret}(K) := \sum_{k=1}^K (V_0^* - V_0^{\pi^k}).$$

Maximum Conditional Total Variance

We define $\text{Var}_{\max}^c$, the maximum conditional total variance.

$$\text{Var}_h^*(s, a) := \mathbb{V}^{r \sim R_{s,a,h}, s' \sim P_{s,a,h}} [r + V_{h+1}^*(s')],$$

$$\text{Var}_{\max}^c := \max_{\pi, s, h} \mathbb{E}^\pi \left[\sum_{h'=1}^h \text{Var}_{h'}^*(s_{h'}, a_{h'}) \mid s_h = s \right].$$

We call

$$\text{Var}_{\max} := \max_{\pi} \mathbb{E}^\pi \left[\sum_{h=1}^H \text{Var}_h^*(s_h, a_h) \right]$$

the maximum unconditional total variance.

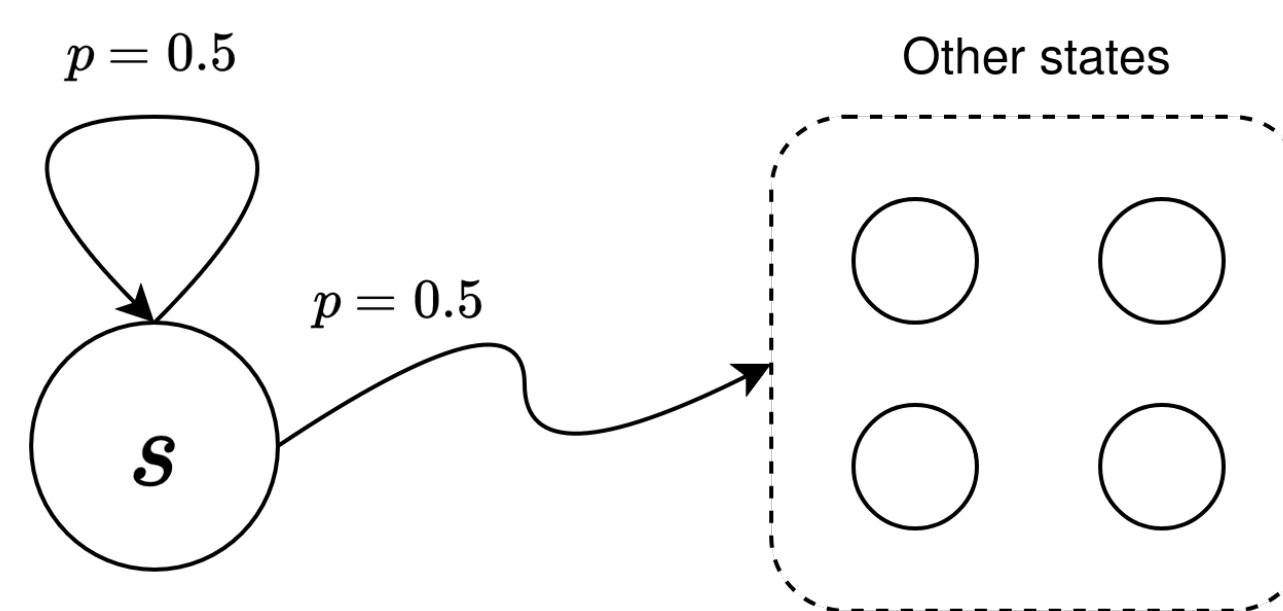


Figure 1. Conditioning on $s_h = s$, $\text{Var}_{\max}^c$ can be as large as $\Theta(H^3)$

Conditions for MDPs

Condition 1: For any possible trajectory, its total reward in a single episode is upper-bounded by H .

Question under Research

We know

$$\mathbb{E}^\pi \left[\sum_{h'=h}^H \text{Var}_{h'}^*(s_{h'}, a_{h'}) \mid s_h = s \right] \leq H^2,$$

but previous gap-dependent analysis does not use this fact.

Can we find the tightest problem-dependent regret while considering both variance and gap?

Regret Bound Comparison

Algorithm	Regret	Variance-Dependent
StrongEuler	$\tilde{O}(\sum_{h,s,a} H Q^* / (\Delta_h(s, a) \vee \Delta_{\min}))$	Yes
	$\tilde{O}(\sum_{(h,s,a) \in \mathcal{Z}_{\text{sub}}} H^3 / \Delta_h(s, a) + \mathcal{Z}_{\text{opt}} H^3 / \Delta_{\min})$	No
Q-Learning (UCB-H)	$\tilde{O}(H^6 S A / \Delta_{\min})$	No
AMB	$\tilde{O}(\sum_{(h,s,a) \in \mathcal{Z}_{\text{sub}}} H^5 / \Delta_h(s, a) + \mathcal{Z}_{\text{mul}} H^5 / \Delta_{\min})$	No
UCB-Advantage	$\tilde{O}((H Q^* + H) H^2 S A / \Delta_{\min})$	No
Q-EarlySettled-Advantage	$\tilde{O}((H Q^* + H) H^2 S A / \Delta_{\min})$	No
MVP This work	$\tilde{O}(\sum_{(h,s,a) \in \mathcal{Z}_{\text{sub}}} W / \Delta_h(s, a) + \mathcal{Z}_{\text{opt}} W / \Delta_{\min})$	Yes

Remark:

- $\tilde{O}$ hides all $\text{poly} \log(SAHK)$ terms, and burn-in terms are not presented in the table. See our paper for more details.
- $W = H^2 \wedge \text{Var}_{\max}^c \leq H Q^* \leq H^3$, so our regret bound recovers all previous results except for AMB.

Lower Bound Result

Given (with some mild conditions):

- Target gap set $\{\Delta_i\}$
- Target maximum conditional total variance L ,

there exists an MDP instance with previous targets so that

$$\text{Regret}(K) \geq \Omega \left(\sum_{i: \Delta_i > 0} \frac{L}{\Delta_i} \cdot \log K \right).$$

Remark:

- Our lower bound matches with the first term of our regret bound.
- The MDP instance can be taken with small maximum unconditional total variance $\text{Var}_{\max}$.