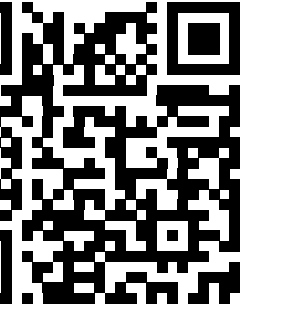


Robust Average-Reward Markov Decision Processes: Minimax-Optimal Learning via Plug-in Reductions

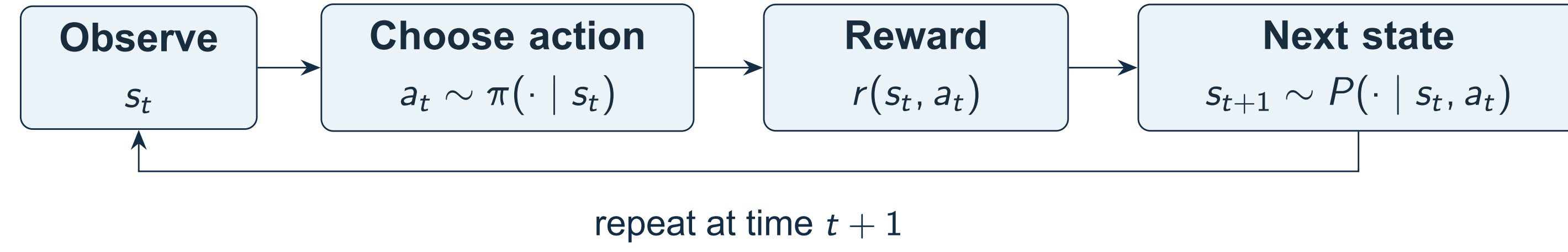
Paper link



Yuepeng Yang (Yale, UPenn) Yuxin Chen (UPenn) Yuejie Chi (Yale)

1 Distributionally robust Markov decision processes

A **Markov decision process** $(\mathcal{S}, \mathcal{A}, P, r)$ consists of a state space $\mathcal{S}$ of size S , an action space $\mathcal{A}$ of size A , a transition kernel P , and a reward function r . A stationary policy $\pi : \mathcal{S} \rightarrow \Delta(\mathcal{A})$ selects actions using the current state:



Uncertainty set and robust objective

For a nominal kernel P^0 and known radii $0 \leq \sigma_{s,a} \leq \sigma$, consider the total-variation uncertainty set with independent perturbations at each state-action pair:

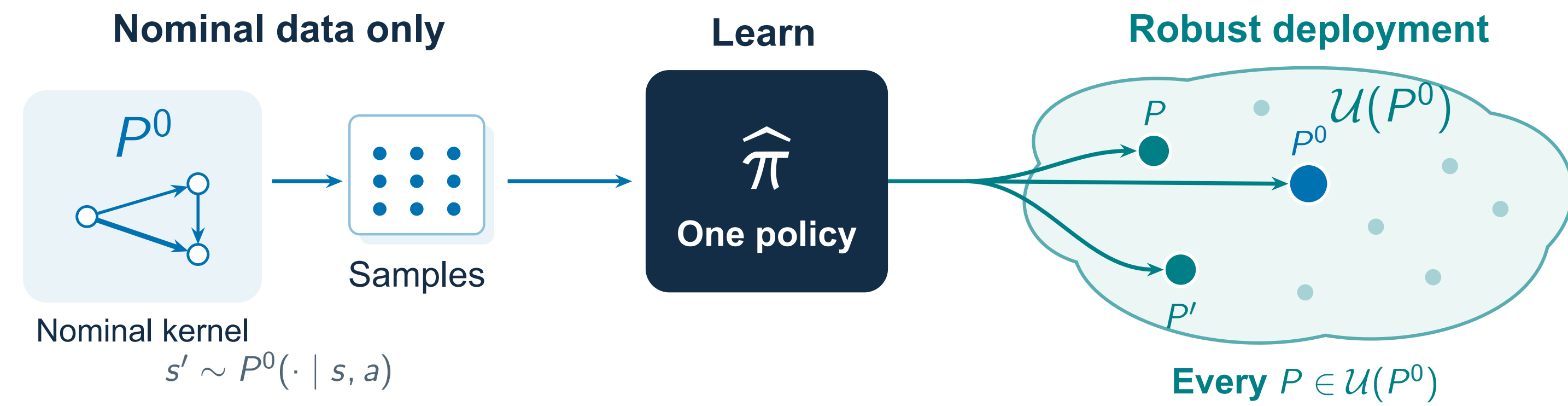
$$\mathcal{U}(P^0) = \prod_{(s,a) \in \mathcal{S} \times \mathcal{A}} \{P_{s,a} \in \Delta(\mathcal{S}) : \|P_{s,a} - P^0_{s,a}\|_{\text{TV}} \leq \sigma_{s,a}\}.$$

We maximize the **worst-case long-run average reward**:

$$\rho^{\pi,\sigma}(s) = \inf_{P \in \mathcal{U}(P^0)} \lim_{T \rightarrow \infty} \mathbb{E}_P^{\pi} \left[\frac{1}{T} \sum_{t=0}^{T-1} r(s_t, a_t) \mid s_0 = s \right].$$

Generative model and sample complexity

- Assume access to a simulator of the nominal transition kernel P^0 .
- Collect N samples per state-action pair, for a total of NSA samples.



For $\varepsilon > 0$ and $\delta \in (0, 1)$, how can we learn $\hat{\pi}$ such that, with probability at least $1 - \delta$,

$$\rho^{*,\sigma}(s) - \rho^{\hat{\pi},\sigma}(s) \leq \varepsilon \quad \text{for every } s \in \mathcal{S}?$$

How many samples are needed?

Difficulty parameters: bias spans

- Starting states share the same average reward, but their early rewards can differ.
- The **nominal** and **robust bias functions**, $h_{P^0}^{\pi}$ and $h^{\pi,\sigma}$, measure the expected cumulative (robust) reward surplus above the long-run average under policy π .
- With $\|h\|_{\text{span}} = \max_s h(s) - \min_s h(s)$, the **nominal** and **robust optimal bias spans** are

$$H_0 = \max \{1, \|h_{P^0}^*\|_{\text{span}}\}, \quad H_{\sigma} = \max \{1, \|h^{*,\sigma}\|_{\text{span}}\}.$$
- These spans measure difficulty, analogous to $(1 - \gamma)^{-1}$ in discounted MDPs.

2 Learning through discounted plug-in reductions

Discounted values approximate average rewards

- For discount factor $0 < \gamma < 1$, the **robust discounted value** is

$$V_{\gamma}^{\pi,\sigma}(s) = \inf_{P \in \mathcal{U}(P^0)} \mathbb{E}_P^{\pi} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s \right].$$

- The optimal value $V_{\gamma}^{*,\sigma}(s) = \max_{\pi} V_{\gamma}^{\pi,\sigma}(s)$ satisfies

$$|(1 - \gamma)V_{\gamma}^{*,\sigma}(s) - \rho^{*,\sigma}| \leq (1 - \gamma)H_{\sigma}.$$

As γ approaches 1, the effective horizon $1/(1 - \gamma)$ increases



- The returned policy $\hat{\pi}$ satisfies

$$\rho^{*,\sigma} - \rho^{\hat{\pi},\sigma} \leq 2(1 - \gamma)H_{\sigma} + (1 - \gamma) \left\| V_{\gamma}^{*,\sigma} - V_{\gamma}^{\hat{\pi},\sigma} \right\|_{\infty}.$$

approximation + discounted policy error

- Ideal choice of γ :** $1 - \gamma \asymp \varepsilon/H_{\sigma}$ gives an approximation error of $O(\varepsilon)$.

Nominal and robust discounted MDP solvers

- Estimate $\hat{P}^0$ from observed transition frequencies and initialize $\hat{V}_{\gamma,0}^{\sigma}(s) = 0$.
- Robust value iteration:**

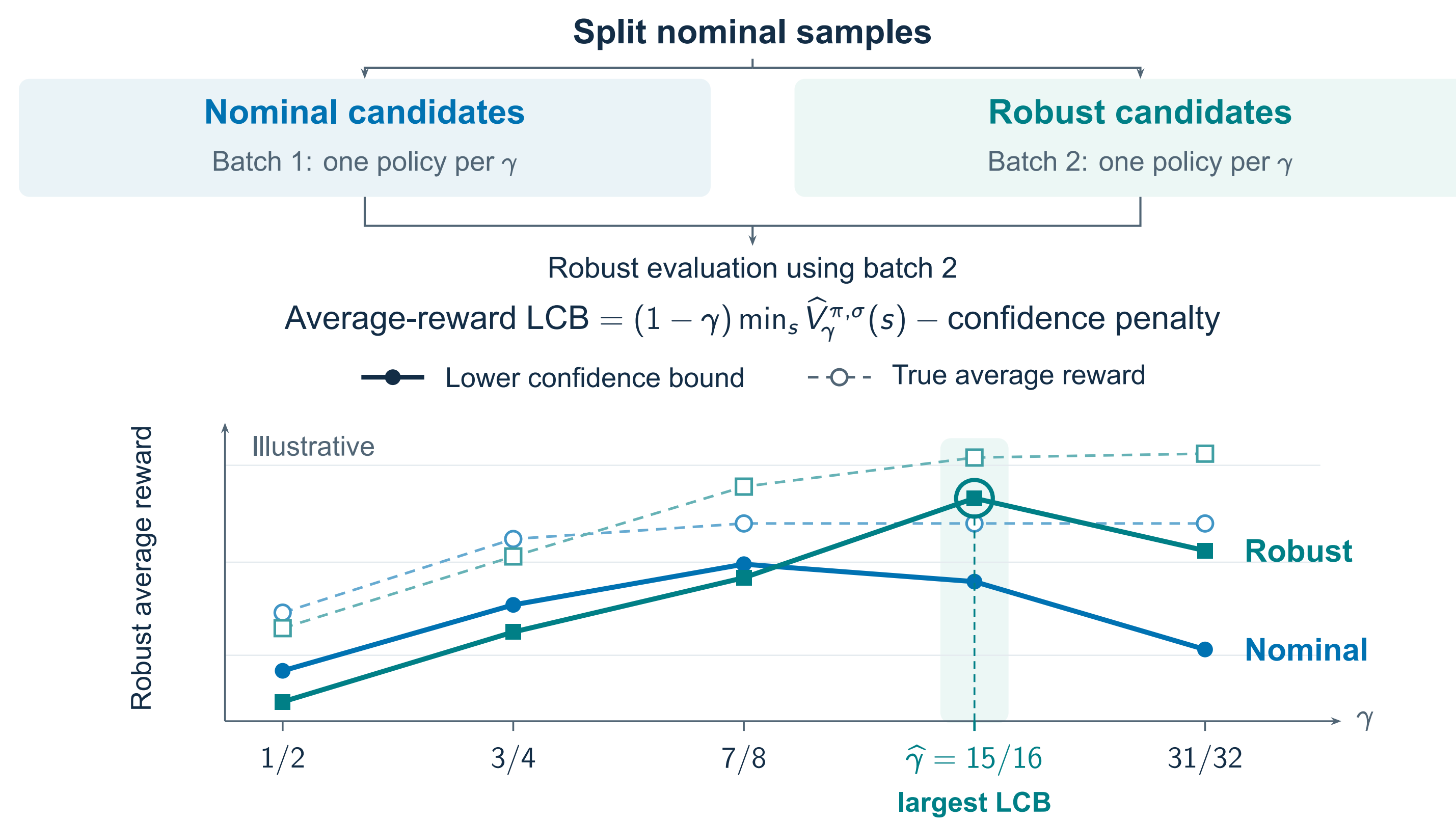
$$\hat{V}_{\gamma,k+1}^{\sigma}(s) \leftarrow \max_{a \in \mathcal{A}} \left\{ r(s, a) + \gamma \min_{q \in \mathcal{U}_{s,a}(\hat{P}^0)} \sum_{s' \in \mathcal{S}} q(s') \hat{V}_{\gamma,k}^{\sigma}(s') \right\}.$$

immediate reward + $\gamma \times$ robust next-state value

- The nominal variant uses $\hat{P}^0$ directly in the Bellman update.

Span-agnostic reduction

When the spans are unknown, try **both reductions** over a finite grid $\gamma = 1 - 2^{-k}$:



Why both reductions? The nominal branch is better when $H_0 < H_{\sigma}$ and $\varepsilon \gtrsim \sigma H_0$.

3 Main result: the statistical price of robustness

Minimax sample complexity rates in two regimes

Regime	Condition	Total samples
High tolerance	$\varepsilon \gtrsim \sigma H_0$	$NSA \asymp \frac{SA \min\{H_0, H_{\sigma}\}}{\varepsilon^2}$
Low tolerance	$\varepsilon \lesssim \sigma H_0$	$NSA \asymp \frac{SA(\min\{H_0, H_{\sigma}\} + \sigma H_{\sigma}^2)}{\varepsilon^2}$

The span-agnostic reductions achieve these rates up to logarithmic factors.

Why the scale σH_0 ?

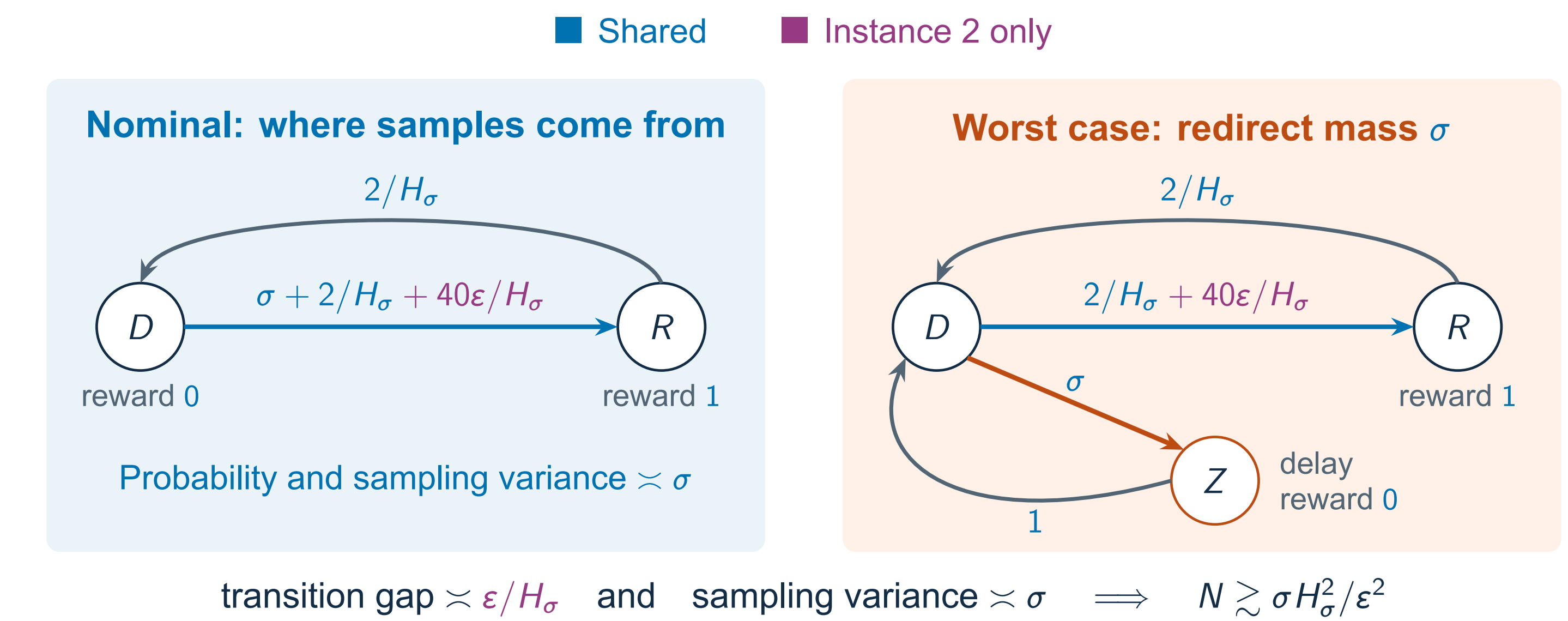
There exists a nominal optimal policy π^* whose robust performance satisfies

$$\rho^{*,\sigma} - \rho^{\pi^*,\sigma} \leq \sigma H_0.$$

A suitable nominal optimal policy therefore meets the robust target when $\varepsilon \gtrsim \sigma H_0$.

Analysis overview

- Span-aware variance analysis.** ($H_{\sigma}^2 \rightarrow H_{\sigma} + \sigma H_{\sigma}^2$)
 - ✗ Bound all value-function spans by $(1 - \gamma)^{-1}$.
 - ✓ Bound a core set of value spans by $O(H_{\sigma})$ plus estimation errors.
- Using a nominal optimal policy as a reference point.** ($H_{\sigma} \rightarrow \min\{H_0, H_{\sigma}\}$)
- Hard instances with amplified robust gaps.**



Adaptive choice of reduction and discount factor

