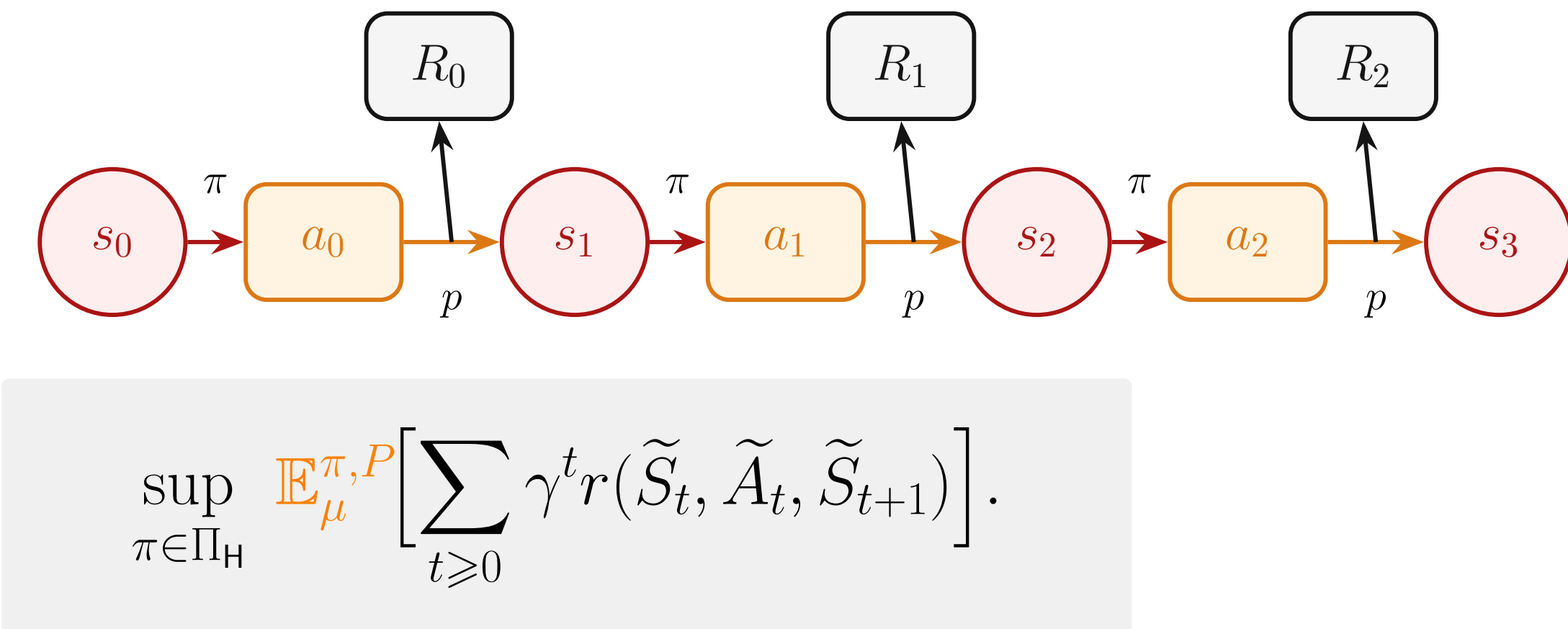




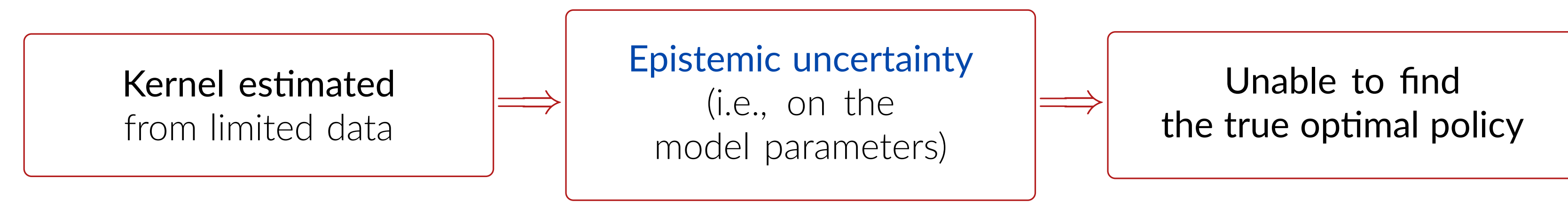
TL;DR: (i) Treating transition probabilities as random variables gives a **unified language for epistemic uncertainty**; (ii) **Bellman Dynamic Programming is only compatible with pessimistic, optimistic and neutral risk-attitudes.**

Model Estimation Errors in MDPs

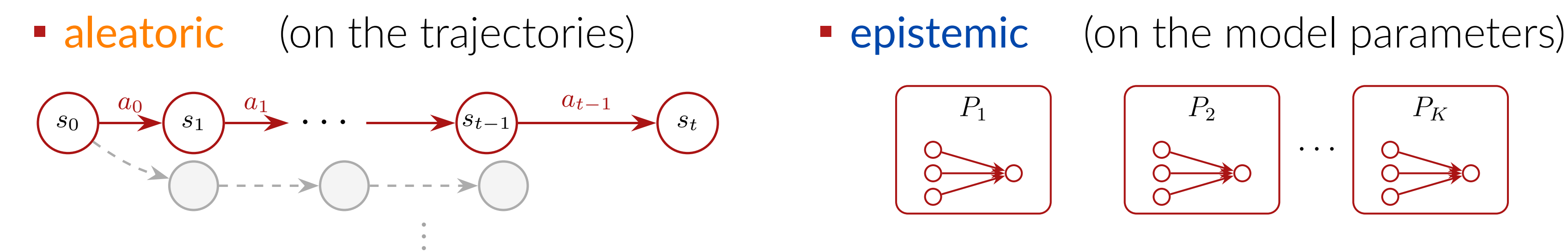
In nominal discounted MDPs $(\mathcal{S}, \mathcal{A}, P, r, \gamma, \mu)$, the transition kernel P is **fixed and known**, and the objective is to maximize the expected return:



In practice, the kernel is estimated from limited data, which can be critical in high-stakes applications:



Two types of uncertainty coexist:



Many existing Models for Epistemic Uncertainty

Model	Treatment of uncertainty	DP Compatible?
Robust MDP	worst plausible transition kernel	yes
Multi-model MDP	average over plausible kernels	sometimes yes
Percentile optimization	optimize a quantile of return	no
CVaR, entropic risk	smoother risk-aware objectives	?

Problem: All these models were introduced independently, with *ad hoc* analyses.

Main Reseach Questions:

Can we unify these approaches? Which epistemic-uncertainty objectives are compatible with Bellman-style dynamic programming?

Ambiguity-Averse MDPs

Core modeling steps:

- Treat the transition kernel as a **random variable**: $\tilde{P} \sim \nu$.
- Policies π induce (random) expected returns $\mu^\top V^{\pi, \tilde{P}} = \mathbb{E}_{\mu, \tilde{P}} (\sum_{t=0}^{\infty} \gamma^t r(\tilde{S}_t, \tilde{A}_t, \tilde{S}_{t+1}))$.
- Use a **risk measure** $\rho : \{\text{RVs}\} \rightarrow \mathbb{R}$ to **evaluate** them and solve $\sup_{\pi \in \Pi_H} \rho(\mu^\top V^{\pi, \tilde{P}})$.

Definition

An **ambiguity-averse MDP** $(\mathcal{S}, \mathcal{A}, \nu, r, \gamma, \mu)$ is an MDP with a random transition kernel $\tilde{P} \sim \nu$, where ν is a distribution over transition kernels.

Objective. For a risk measure $\rho : \{\text{RVs}\} \rightarrow \mathbb{R}$, find

$$\pi^* \in \arg \sup_{\pi \in \Pi_H} \rho(\mu^\top V^{\pi, \tilde{P}}).$$

Ambiguity-averse MDPs recover most epistemic-uncertainty models in the litterature:

- $\nu = \delta_P$: nominal MDP.
- $\rho = \mathbb{E}^\nu$: multi-model MDP.
- $\rho = \text{ess inf}$: robust MDP.
- $\rho = \text{ess sup}$: optimistic MDP.
- $\rho = \text{VaR}_\alpha$: percentile objective.
- $\rho = \text{CVaR}_\alpha$: tail-risk objective.
- $\rho = \dots$: any other risk attitude.

Two ways Epistemic Uncertainty Evolves

Static kernels. One kernel is sampled from ν at $t = 0$ and remains fixed:

$$\tilde{P}_0 = \tilde{P}_1 = \tilde{P}_2 = \dots$$

Resampled kernels. A fresh plausible kernel is sampled from ν at each period $t \in \mathbb{N}$:

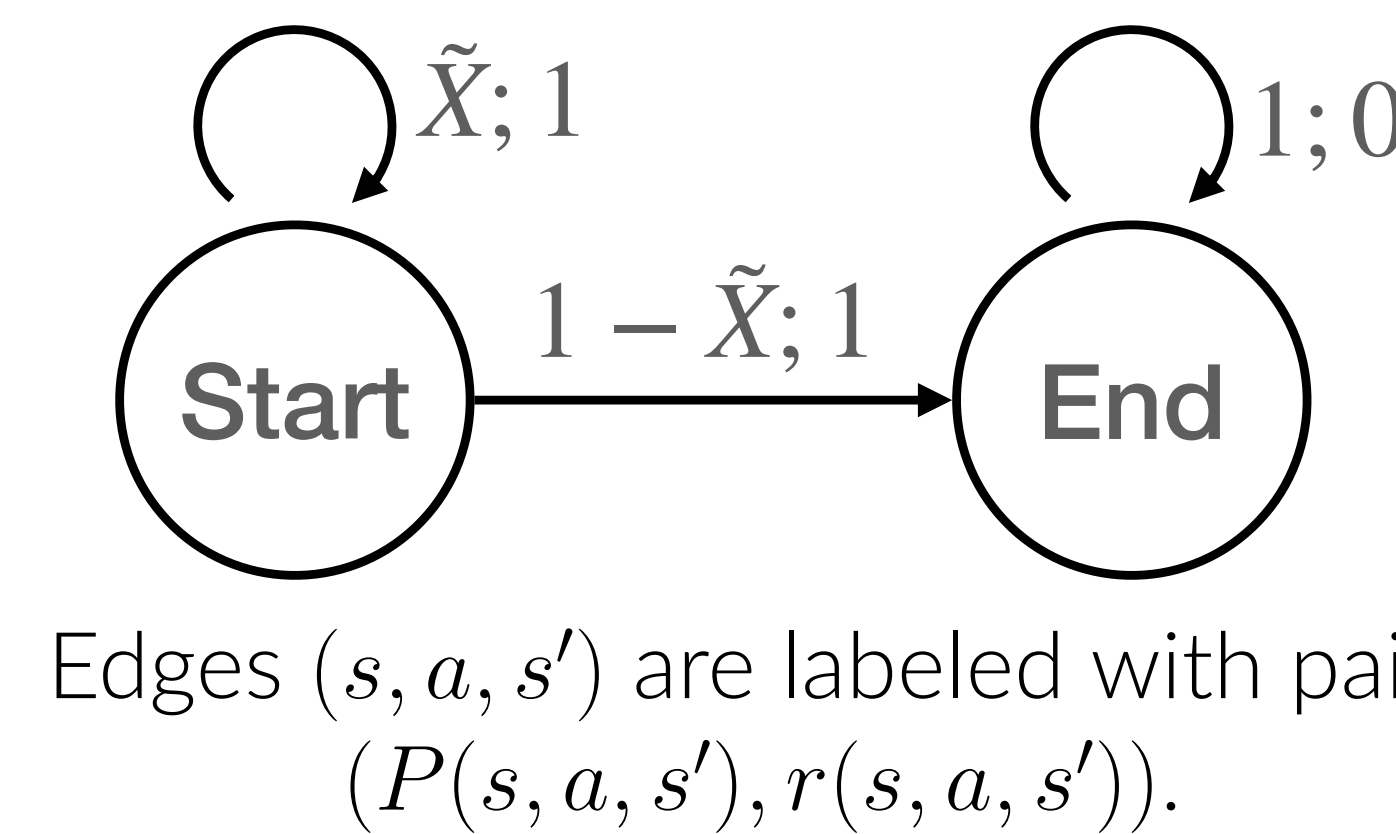
$$\tilde{P}_0, \tilde{P}_1, \tilde{P}_2, \dots \text{ i.i.d.}$$

A Static Kernel Example

MDP with two states, a **unique action** and $\tilde{X} \sim \text{Unif}([0, 1])$.

For $\rho = \mathbb{E}^\nu$:

$$\mathbb{E} \left[V_{\text{Start}}^{\pi, \tilde{P}} \right] = \mathbb{E} \left[\frac{1}{1 - \gamma \tilde{X}} \right] = \frac{-\log(1 - \gamma)}{\gamma}.$$



Ambiguity-Averse Dynamic Programming

For each state s , define the **ambiguity-averse value** of a policy and the **optimal value**:

$$V^{\pi, \nu, \rho}(s) = \rho \left(V^{\pi, \tilde{P}}(s) \right), \quad V^{\nu, \rho}(s) = \sup_{\pi \in \Pi_H} V^{\pi, \nu, \rho}(s).$$

For a stationary policy $\pi \in \Pi_S$, define the **ambiguity-averse Bellman operator** and the **optimal operator**:

$$T^{\pi, \nu, \rho} v(s) = \rho \left(T^{\pi, \tilde{P}} v(s) \right), \quad T^{\nu, \rho} v(s) = \sup_{\pi \in \Pi_S} T^{\pi, \nu, \rho} v(s).$$

We say that ρ satisfies the **Dynamic Programming principles** if the value functions are fixed points of the Bellman operators:

$$V^{\pi, \nu, \rho} = T^{\pi, \nu, \rho} V^{\pi, \nu, \rho}, \quad V^{\nu, \rho} = T^{\nu, \rho} V^{\nu, \rho}.$$

Proposition. If ρ satisfies the Dynamic Programming principles, then

- there exists a stationary optimal policy $\pi^* \in \Pi_S$.
- value and policy iteration algorithms inherit the usual convergence guarantees.

Main Results

Proposition. If ρ is monotone, law-invariant and satisfies the DP principles, then

- $\rho(\tilde{X} + \tilde{Y}) = \rho(\tilde{X}) + \rho(\tilde{Y})$ for $\tilde{X} \perp\!\!\!\perp \tilde{Y}$.
- $\rho(\alpha \tilde{X}) = \alpha \rho(\tilde{X})$ for $\alpha \geq 0$.

Theorem. A monotone, law-invariant ρ satisfies the DP principles if and **only if**

- $\rho \in \{\text{ess inf}, \text{ess sup}\}$ (Static kernel case).
- $\rho \in \{\text{ess inf}, \text{ess sup}, \mathbb{E}\}$ (Resampled kernel case).

Takeaways

- Model uncertainty leads to misspecified policies.
- Ambiguity-Averse MDPs: Unified framework for epistemic uncertainty.
- A full characterization for risk-attitudes compatible with DP.
- Need to go beyond classical DP to use richer risk-attitudes.