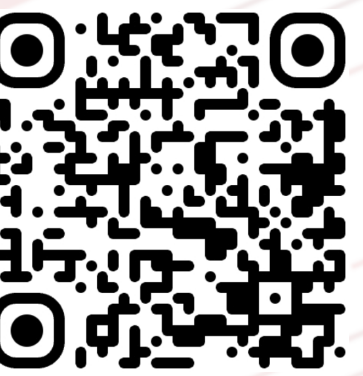


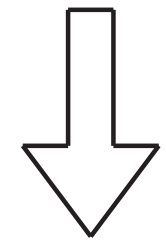
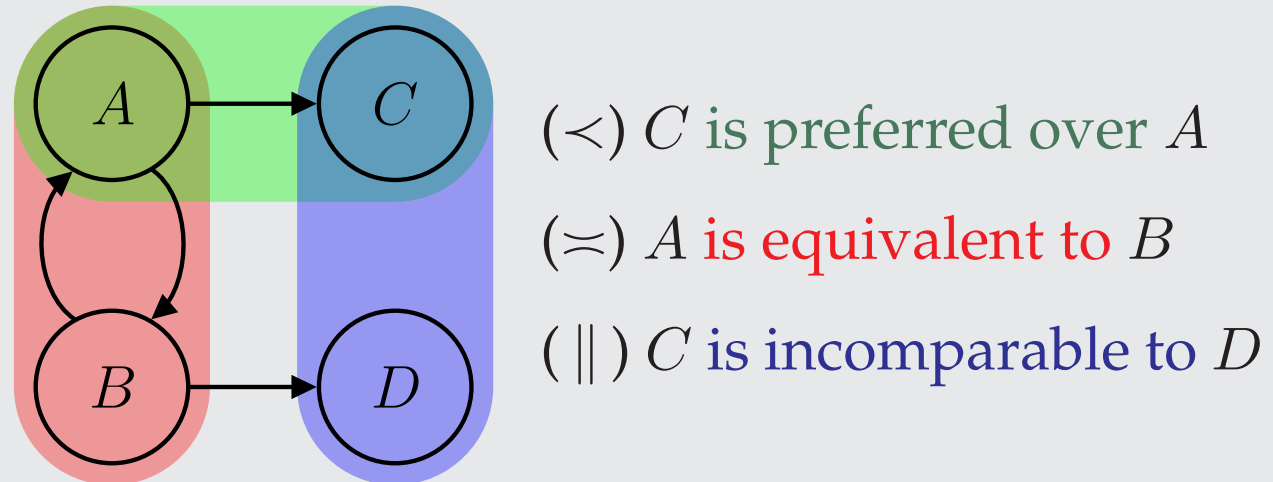


SETTING

PREFERENCE

$$\text{PrefMDP} = (\underbrace{\mathcal{S}, \mathcal{A}, H, p, \mu}_{\text{MDP} \setminus \mathcal{R}}, \leq_{\mathcal{T}})$$

$\leq_{\mathcal{T}} \subseteq \mathcal{T} \times \mathcal{T}$ is a **partial (pre)order** over the trajectory space $\mathcal{T}$



Assumption: Human expresses preferences based on an underlying (unknown) utility

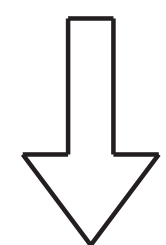
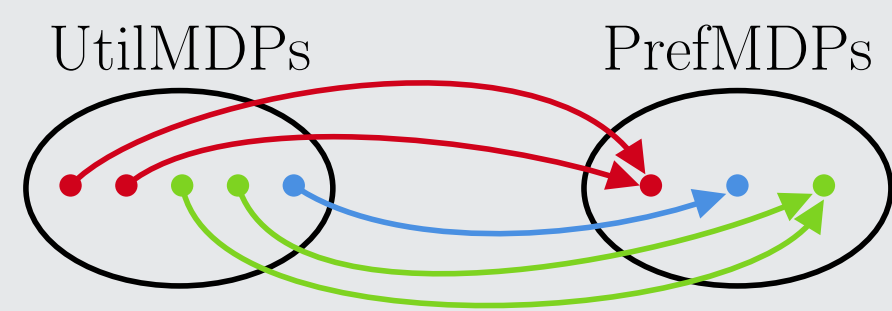
$$\text{UtilMDP} = (\underbrace{\mathcal{S}, \mathcal{A}, H, p, \mu}_{\text{MDP} \setminus \mathcal{R}}, \mathbf{u})$$

$\mathbf{u} : \mathcal{T} \rightarrow \mathbb{R}^m$ is a **multi-dimensional utility**

Expected utility of a policy $\pi \in \Pi$:

$$J(\pi; \mathbf{u}) := \sum_{\tau \in \mathcal{T}} d_{\pi}(\tau) \mathbf{u}(\tau) = \langle d_{\pi}, \mathbf{u} \rangle,$$

where d_{π} is the distribution over trajectories induced by π



Assumption: Human expresses preferences based on an (unknown) Markovian reward

$$\text{MDP} = (\underbrace{\mathcal{S}, \mathcal{A}, H, p, \mu}_{\text{MDP} \setminus \mathcal{R}}, \mathbf{r})$$

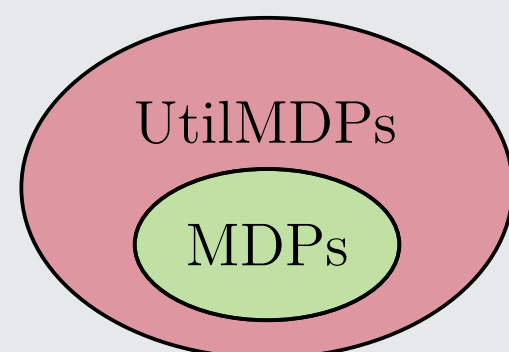
$\mathbf{r} := (r_h)_{h \in [H]}$ is a **stage-dependent multi-dimensional reward function**

Trajectory return:

$$\mathbf{u}_{\mathbf{r}}(\tau) := \sum_{h=1}^H r_h(s_h, a_h)$$

Expected policy return:

$$J(\pi; \mathbf{r}) := J(\pi; \mathbf{u}_{\mathbf{r}})$$



REWARD

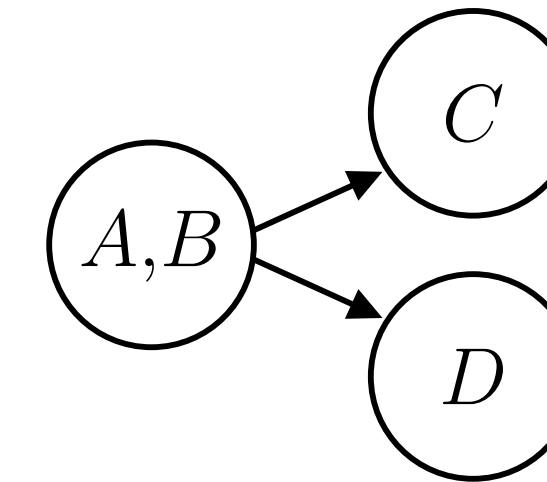
COMPATIBLE UTILITIES

COMPATIBLE UTILITY

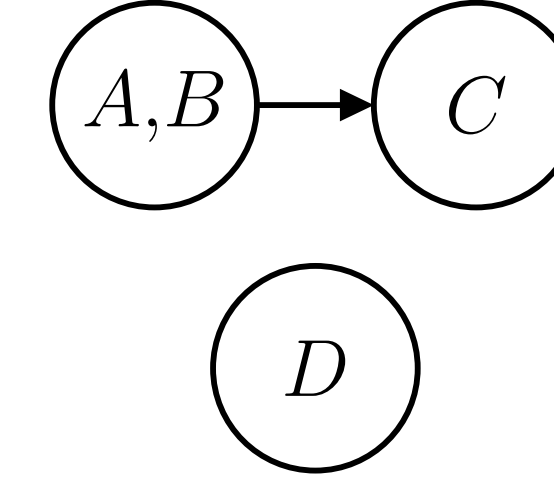
$\mathbf{u}$ is **compatible** with $\leq_{\mathcal{T}}$ if $\forall \tau, \tau' \in \mathcal{T}$:
 $\tau \leq_{\mathcal{T}} \tau' \Rightarrow \mathbf{u}(\tau) \leq \mathbf{u}(\tau')$ (element-wise)

Realizer Dimension	Exists?	Computational Complexity
$< \dim(\leq_{\mathcal{T}})$	✗	—
$= \dim(\leq_{\mathcal{T}})$	✓	NP-hard
$> \dim(\leq_{\mathcal{T}})$	✓	$\text{Poly}(\mathcal{T})$

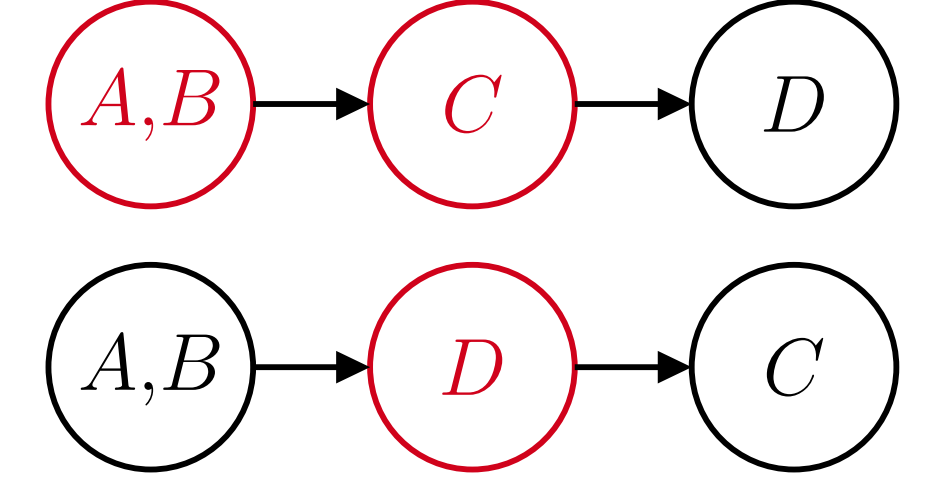
where $\dim(\cdot)$ is the order dimension



(i) Represent the quotient set w.r.t. equivalence of $\leq_{\mathcal{T}}$ as a DAG



(ii) Solve a *minimum path cover* problem to obtain a set of $\text{width}(\leq_{\mathcal{T}})$ chains



(iii) Extend each chain to obtain a linear extension of $\leq_{\mathcal{T}}$, accounting for incomparabilities

Procedure to derive a realizer of size $\text{width}(\leq_{\mathcal{T}}) \geq \dim(\leq_{\mathcal{T}})$ in $\mathcal{O}(|\mathcal{T}|^3)$

POLICY DOMINANCE

POLICY DOMINANCE

Policy π **$\leq_{\mathcal{T}}$ -strictly dominates** policy π' ($\pi' <_{\Pi} \pi$) if it yields a strictly better expected utility

$$J(\pi; \mathbf{u}) - J(\pi'; \mathbf{u}) > 0 \quad (\text{element-wise})$$

for every compatible utility function $\mathbf{u}$

$\leq_{\mathcal{T}}$ -PARETO OPTIMALITY

Set of $\leq_{\mathcal{T}}$ -Pareto optimal policies:

$$\Pi^*(\leq_{\mathcal{T}}) := \{\pi \in \Pi : \nexists \pi' \in \Pi \text{ s.t. } \pi <_{\Pi} \pi'\}$$

HOW TO EVALUATE?

$\pi' \leq_{\Pi} \pi$ can be verified by evaluating if:

$$\forall n \in [|\mathcal{T}|] : \sum_{i=1}^n (d_{\pi}(i) - d_{\pi'}(i)) \geq 0$$

holds for every **linear extension** of $\leq_{\mathcal{T}}$, where index i represents the i -th trajectory sorted w.r.t. the linear extension

OPEN QUESTION

Is there an efficient evaluation method?

Trivial solution requires the evaluation of $\mathcal{O}(|\mathcal{T}|!)$ linear extensions of $\leq_{\mathcal{T}}$

APPROXIMATION VIA MARKOVIAN REWARDS

WHY THE NEED FOR APPROXIMATION?

Preferences

- + Most general type of feedback
- Intractable without introducing a structure

Utilities

- + Assign numerical signals to each trajectory
- Complexity of learning is exponential in H

Rewards

- + Enable efficient learning
- Less representational power

Human-friendliness ↑
Tractability ↓

(CONVEX) QUADRATIC PROGRAM

Goal: Find $\mathbf{u}$ and $\mathbf{r}$ that best represent $\leq_{\mathcal{T}}$

Input: A realizer $\{\leq_{\mathcal{T}, i}\}_{i \in [m]}$ of $\leq_{\mathcal{T}}$ of size m

Idea: Jointly choose $\mathbf{u}$ compatible with the realizer and $\mathbf{r}$ as Markovian approximation of $\mathbf{u}$

Define: $B \in \{0, 1\}^{|\mathcal{T}| \times |S||A|^H}$ (binary encoding of $\mathcal{T}$) and $A := I_{|\mathcal{T}|} - B(B^{\top}B)^{-1}B^{\top}$ (OLS)

$$\eta^* := \min \|\mathbf{A}\mathbf{u}\|_F^2$$

$$\text{s.t. } u_j(i+1) \leq u_j(i) - \varepsilon \quad \forall i \in [|\mathcal{T}| - 1], j \in [m]$$

If $\eta^* = 0 \rightarrow \mathbf{u}_{\mathbf{r}} = \mathbf{u} \rightarrow$ Preferences derive from a Markovian reward function

If $\eta^* > 0 \rightarrow$ Preferences cannot be expressed via Markovian rewards

$\rightarrow$ Using $\mathbf{r}$, we learn how to solve a **simpler surrogate** problem

$\rightarrow$ Such an **approximation introduces a suboptimality** in terms of the performance of the optimal policy bounded by $2\sqrt{m\eta^*}$