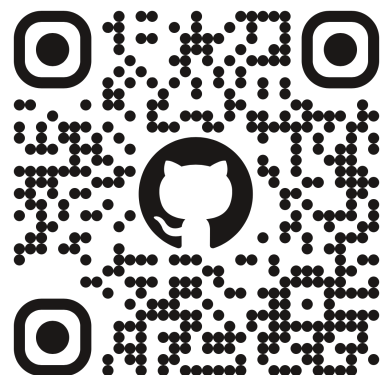


Trust Region Reward Optimization and Proximal Inverse Reward Optimization Algorithm

YANG CHEN

NeurIPS 2025



REINFORCEMENT LEARNING

Human preferences exist

Aristotle, Plato, Epicurus...



PHILOSOPHY & LOGIC

Early

1662



Antoine Arnauld, Pierre Nicole

Port Royal Logic

Quantification of preferences

1823

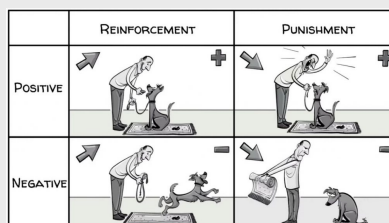


Jeremy Bentham

Utility Calculus

Increasing rewards

Operant Conditioning



1938

1947



Von Neumann, Morgenstern

Utility Theorem

Time and context-independent preferences

Pairwise preferences

Bradley-Terry Model

$$\begin{aligned} \text{logit}(\text{Prob}(i \text{ beats } j)) \\ &= \log \left(\frac{\text{Prob}(i \text{ beats } j)}{\text{Prob}(j \text{ beats } i)} \right) \\ &= \log \left(\frac{p_i}{p_j} \right) = \beta_i - \beta_j. \end{aligned}$$

1952

1957

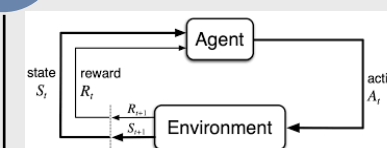


Bellman Equation

Dynamic Programming Equation

Optimal policies exist

MATH. MODELS & THEORIES



Markov Decision Process

LLMs

2025

2022



DRPO



ChatGPT

RLHF

Reinforcement learning from human preferences from data

Deep reinforcement learning

BEAT HUMAN

Alpha Fold 1

DeepMind



2018

2017

PPO Algorithm

OpenAI

$$L^{CLIP}(\theta) = \mathbb{E}_t[\min(r_t(\theta))\hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon)\hat{A}_t)]$$

Alpha Go

DeepMind

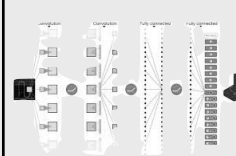


2016

2013

Deep Q learning

DeepMind



1995

1992

1988



IBM TD Gammon

Computing Optimal policies with approximation methods

Q learning

	→	←	↑	↓
Start	0	1	0	0
Idle	2	0	0	3
Hole	0	2	0	0
End	1	0	0	0

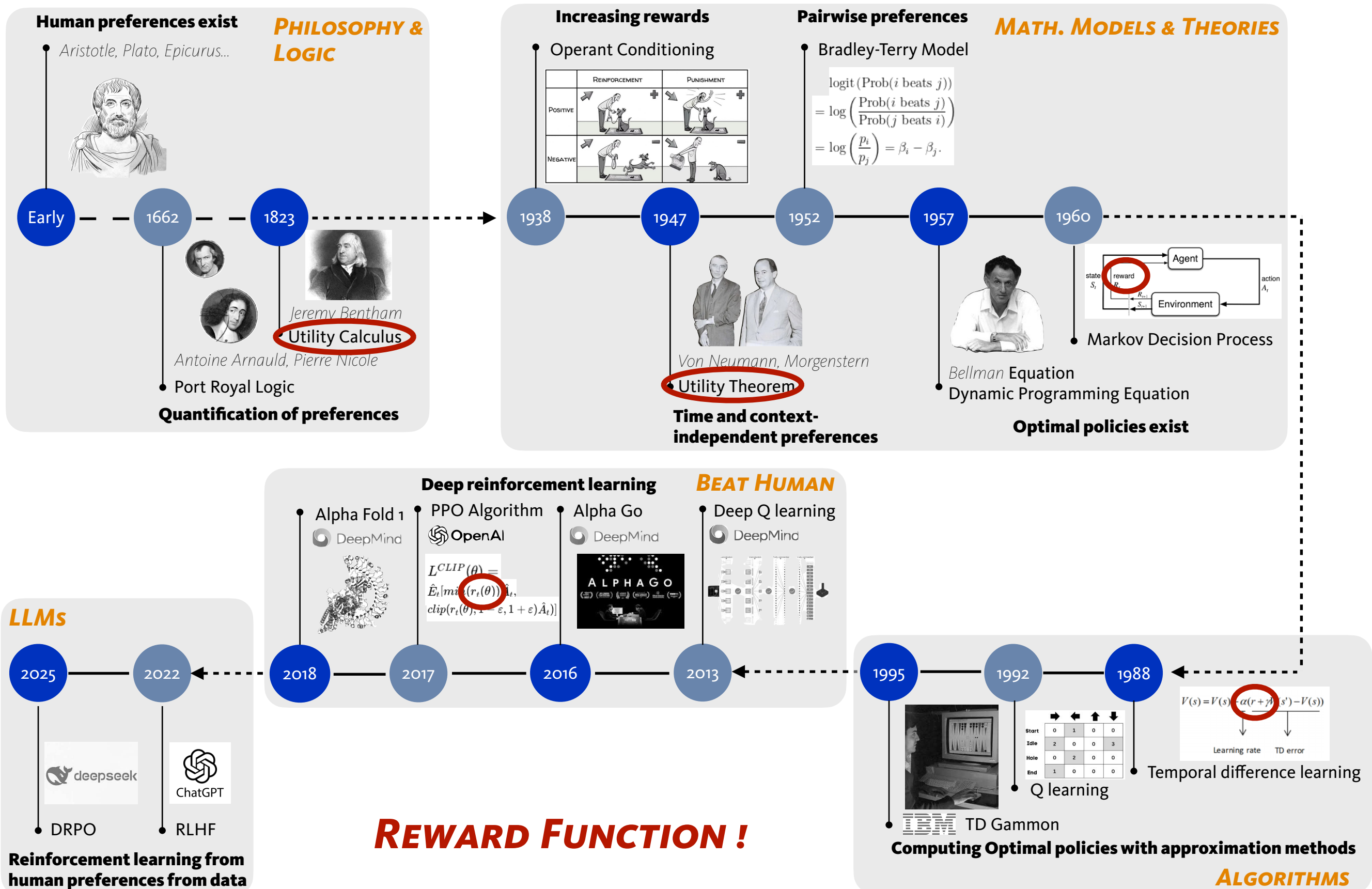
$$V(s) = V(s) + \alpha(r + \gamma V(s') - V(s))$$

Learning rate TD error

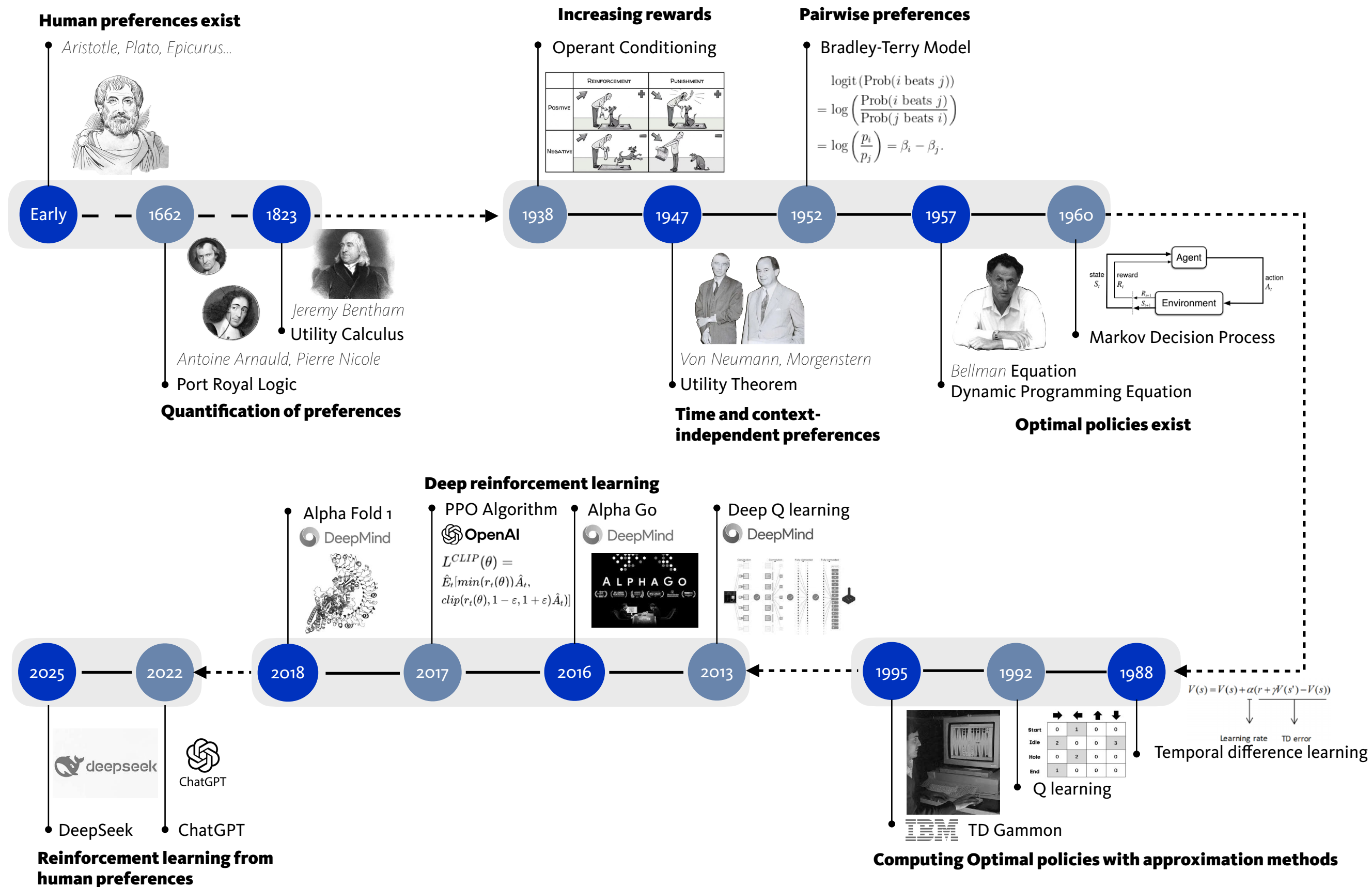
Temporal difference learning

ALGORITHMS

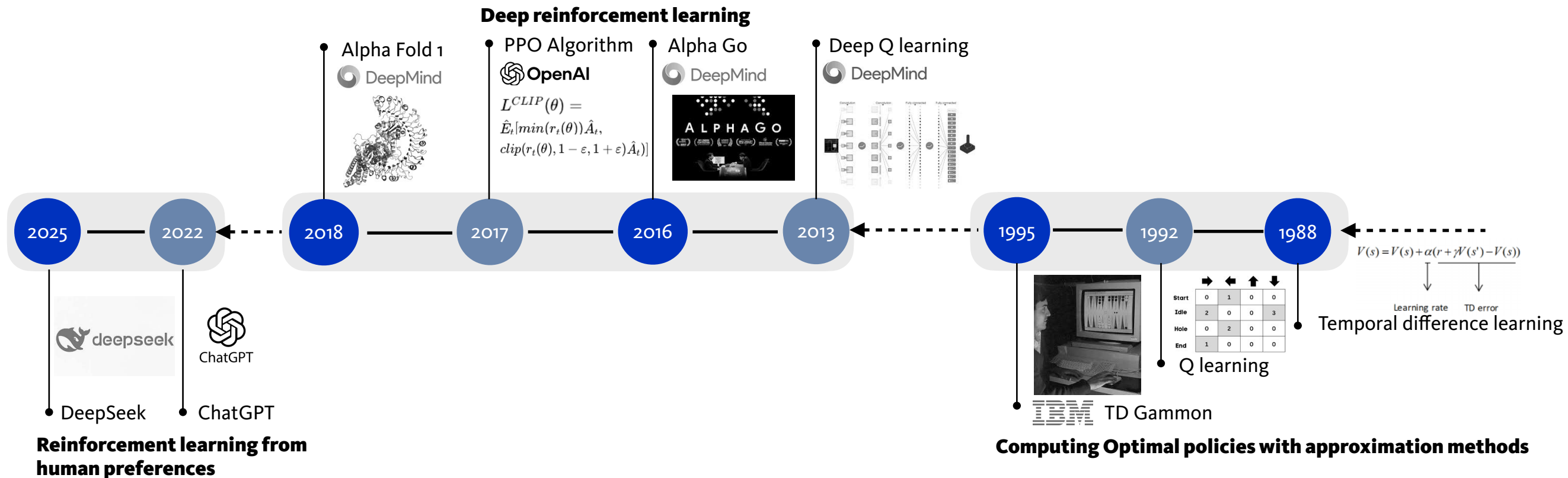
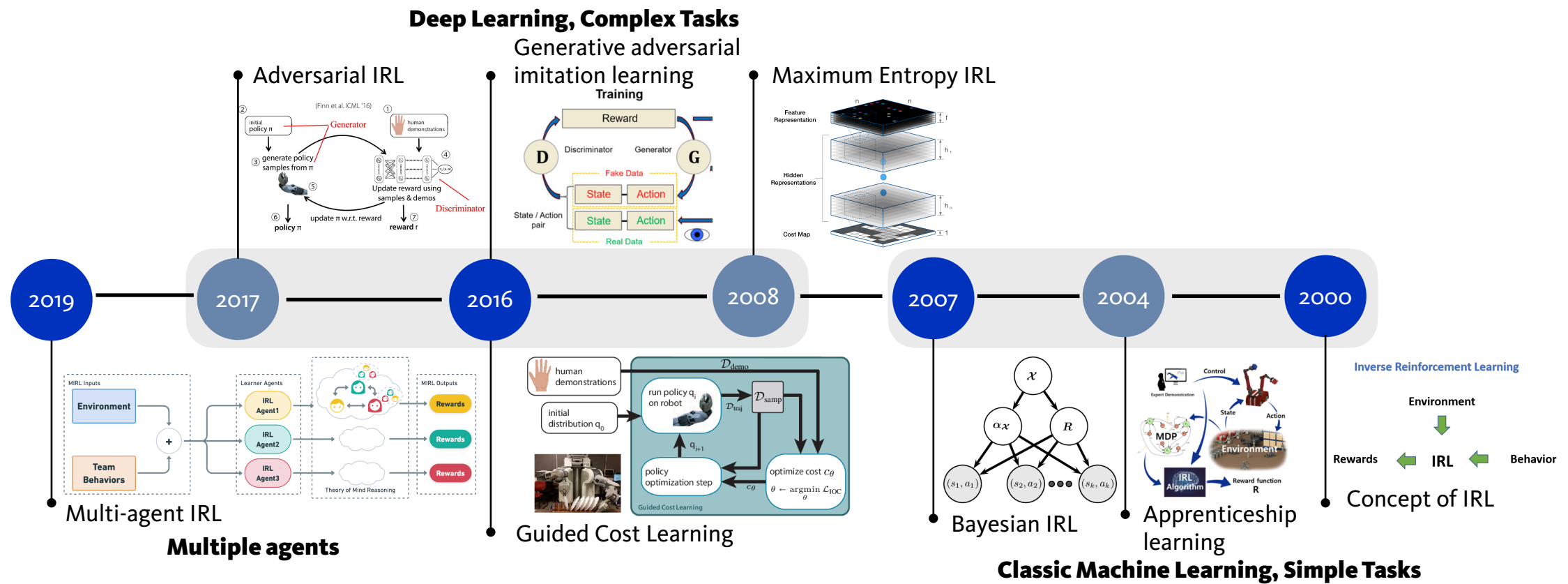
REINFORCEMENT LEARNING



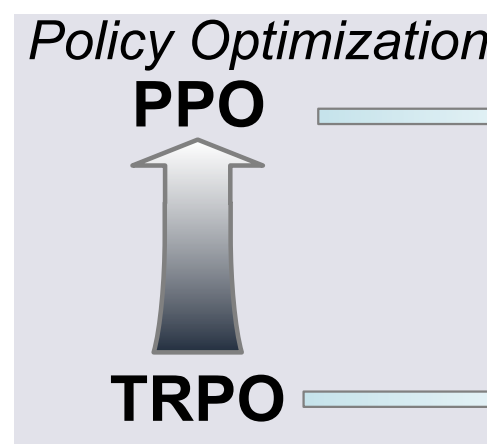
IMITATION LEARNING AND INVERSE REINFORCEMENT LEARNING



IMITATION LEARNING AND INVERSE REINFORCEMENT LEARNING



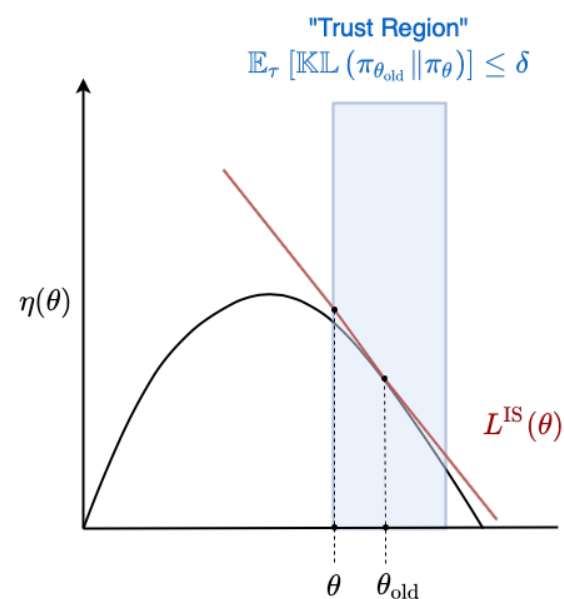
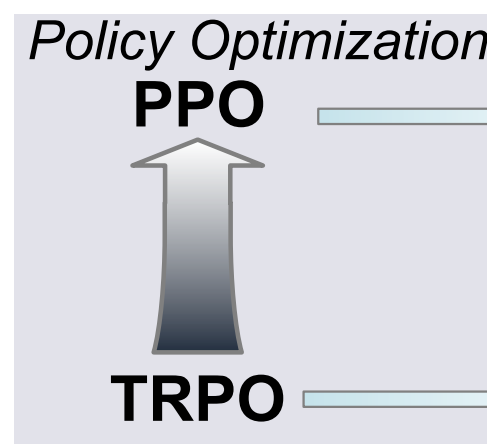
THE SUCCESS OF STABLE REINFORCEMENT LEARNING



Sources of pictures: <https://avandekleut.github.io/assets/ppo/trust-region.png> https://huggingface.co/blog/assets/93_deep_rl_ppo/cliff.jpg

THE SUCCESS OF STABLE REINFORCEMENT LEARNING

$$\begin{aligned} & \underset{\theta}{\text{maximize}} && \hat{\mathbb{E}}_t \left[\frac{\pi_{\theta}(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)} \hat{A}_t \right] \\ & \text{subject to} && \hat{\mathbb{E}}_t [\text{KL}[\pi_{\theta_{\text{old}}}(\cdot | s_t), \pi_{\theta}(\cdot | s_t)]] \leq \delta. \end{aligned}$$



Sources of pictures: <https://avandekleut.github.io/assets/ppo/trust-region.png> https://huggingface.co/blog/assets/93_deep_rl_ppo/cliff.jpg

THE SUCCESS OF STABLE REINFORCEMENT LEARNING

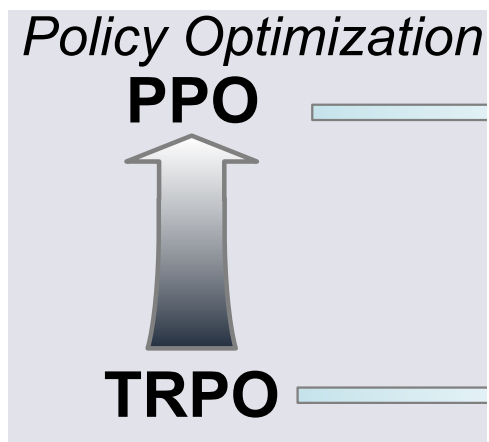
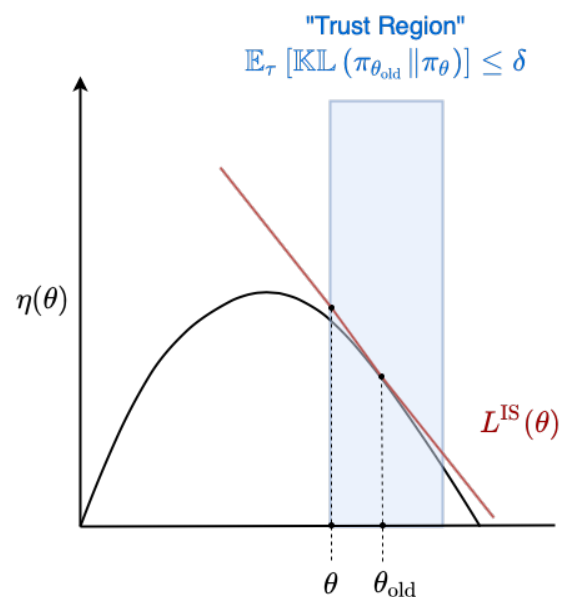
PPO's Clipped surrogate objective function

$$L^{CLIP}(\theta) = \hat{\mathbb{E}}_t \left[\min(r_t(\theta) \hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_t) \right]$$

The ratio function

$$r_t(\theta) = \frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)}$$

$$\begin{aligned} & \underset{\theta}{\text{maximize}} \quad \hat{\mathbb{E}}_t \left[\frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)} \hat{A}_t \right] \\ & \text{subject to} \quad \hat{\mathbb{E}}_t [\text{KL}[\pi_{\theta_{\text{old}}}(\cdot | s_t), \pi_\theta(\cdot | s_t)]] \leq \delta. \end{aligned}$$



Sources of pictures: <https://avandekleut.github.io/assets/ppo/trust-region.png> https://huggingface.co/blog/assets/93_deep_rl_ppo/cliff.jpg

THE SUCCESS OF STABLE REINFORCEMENT LEARNING

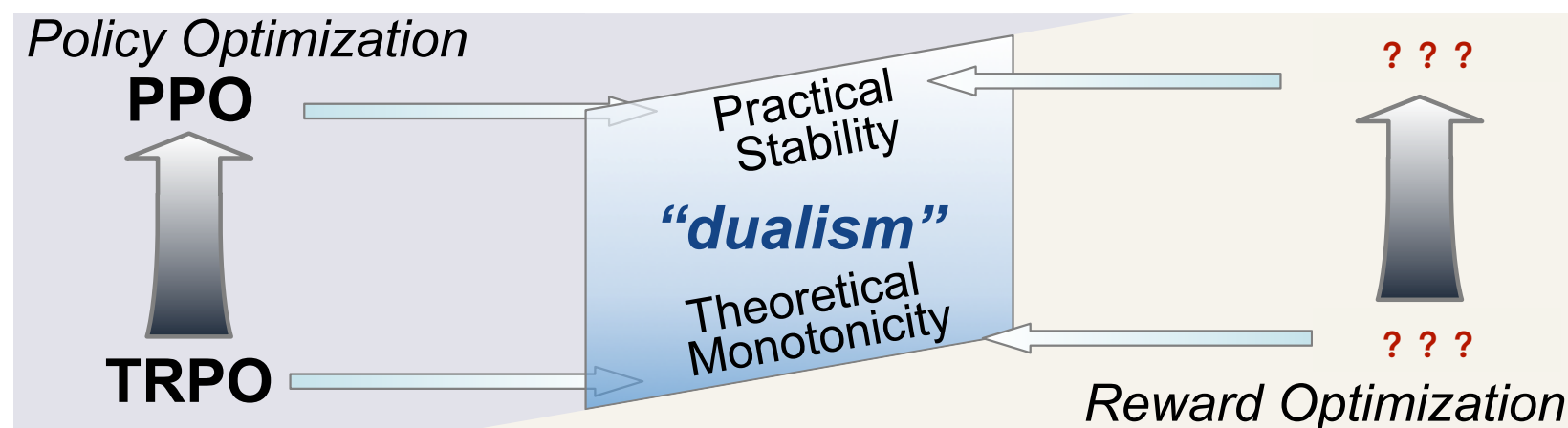
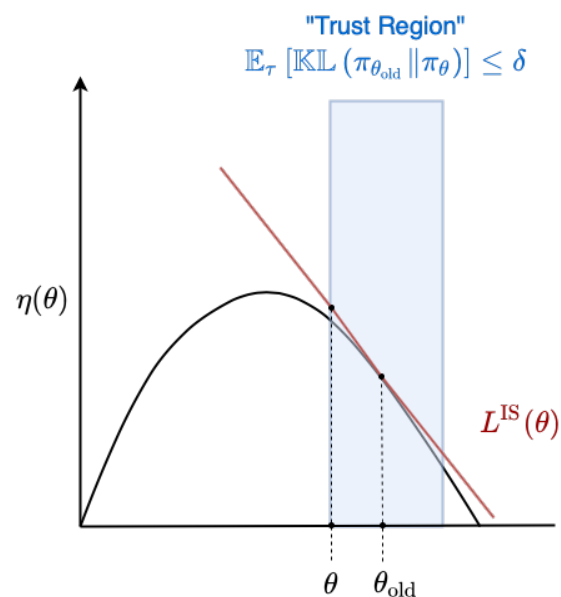
PPO's Clipped surrogate objective function

$$L^{CLIP}(\theta) = \hat{\mathbb{E}}_t \left[\min(r_t(\theta) \hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_t) \right]$$

The ratio function

$$r_t(\theta) = \frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)}$$

$$\begin{aligned} & \underset{\theta}{\text{maximize}} \quad \hat{\mathbb{E}}_t \left[\frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)} \hat{A}_t \right] \\ & \text{subject to} \quad \hat{\mathbb{E}}_t [\text{KL}[\pi_{\theta_{\text{old}}}(\cdot | s_t), \pi_\theta(\cdot | s_t)]] \leq \delta. \end{aligned}$$



Sources of pictures: <https://avandekleut.github.io/assets/ppo/trust-region.png> https://huggingface.co/blog/assets/93_deep_rl_ppo/cliff.jpg

THE SUCCESS OF STABLE REINFORCEMENT LEARNING

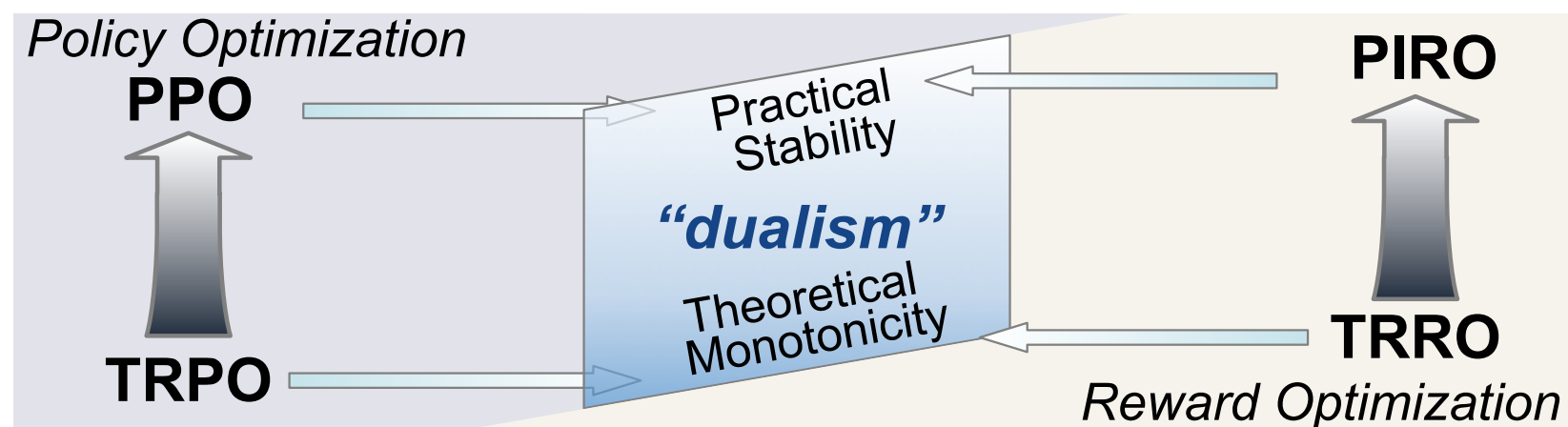
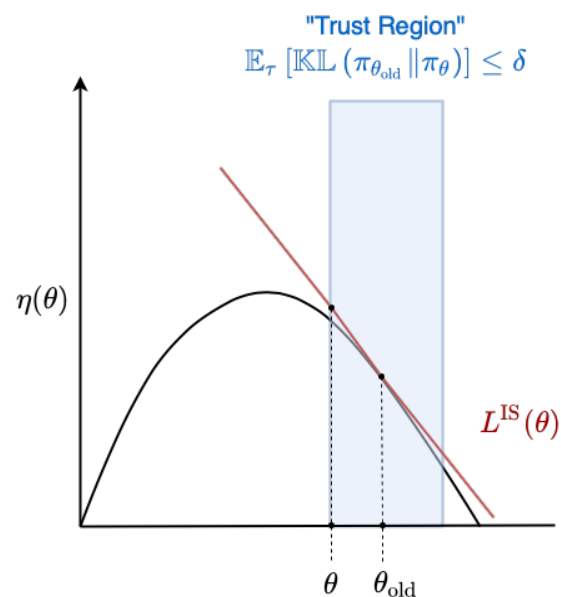
PPO's Clipped surrogate objective function

$$L^{CLIP}(\theta) = \hat{\mathbb{E}}_t \left[\min(r_t(\theta) \hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_t) \right]$$

The ratio function

$$r_t(\theta) = \frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)}$$

$$\begin{aligned} & \underset{\theta}{\text{maximize}} \quad \hat{\mathbb{E}}_t \left[\frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)} \hat{A}_t \right] \\ & \text{subject to} \quad \hat{\mathbb{E}}_t [\text{KL}[\pi_{\theta_{\text{old}}}(\cdot | s_t), \pi_\theta(\cdot | s_t)]] \leq \delta. \end{aligned}$$



Sources of pictures: <https://avandekleut.github.io/assets/ppo/trust-region.png> https://huggingface.co/blog/assets/93_deep_rl_ppo/cliff.jpg

THE SUCCESS OF STABLE REINFORCEMENT LEARNING

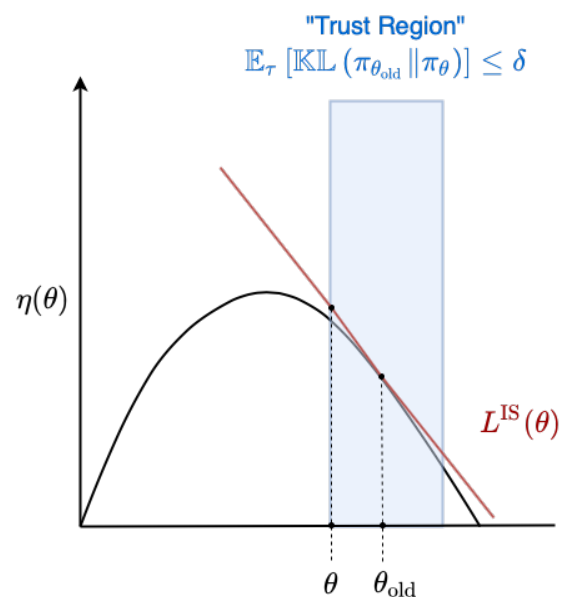
PPO's Clipped surrogate objective function

$$L^{CLIP}(\theta) = \hat{\mathbb{E}}_t \left[\min(r_t(\theta) \hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_t) \right]$$

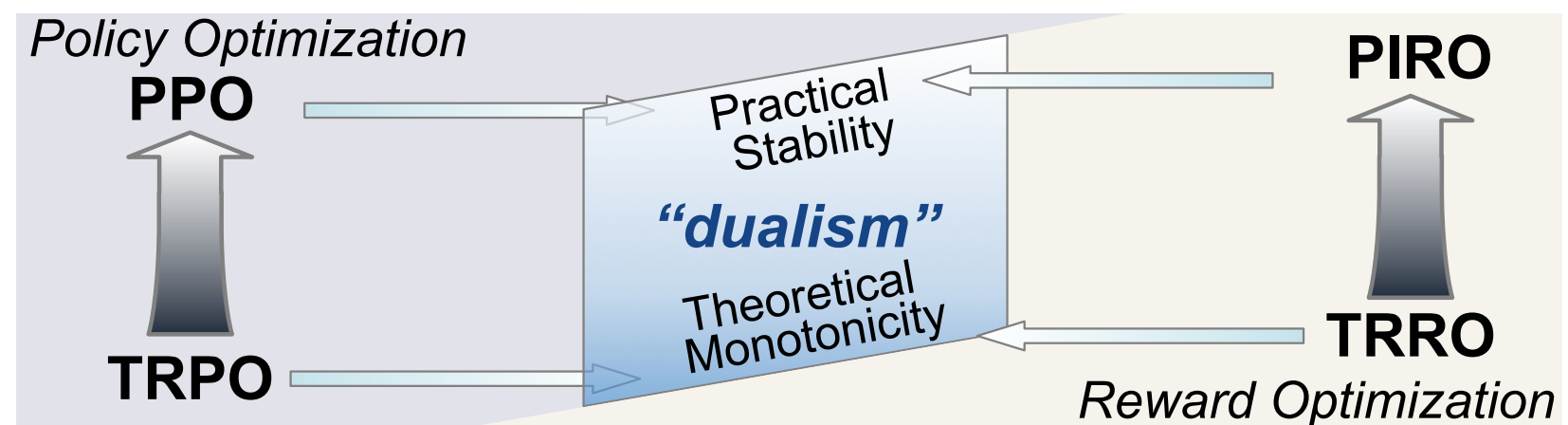
The ratio function

$$r_t(\theta) = \frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)}$$

$$\begin{aligned} & \underset{\theta}{\text{maximize}} \quad \hat{\mathbb{E}}_t \left[\frac{\pi_\theta(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)} \hat{A}_t \right] \\ & \text{subject to} \quad \hat{\mathbb{E}}_t [\text{KL}[\pi_{\theta_{\text{old}}}(\cdot | s_t), \pi_\theta(\cdot | s_t)]] \leq \delta. \end{aligned}$$



Proximal Inverse Reward Optimization



Trust Region Reward Optimization

Sources of pictures: <https://avandekleut.github.io/assets/ppo/trust-region.png> https://huggingface.co/blog/assets/93_deep_rl_ppo/cliff.jpg

CONCEPTS OF INVERSE REINFORCEMENT LEARNING

ELEMENTS OF INVERSE REINFORCEMENT LEARNING

$$\min_{\pi \in \Pi} \max_{r \in \mathcal{R}} J(\pi_E, r) - J(\pi, r) = \mathbb{E}_{\rho^{\pi_E}} [r(\mathbf{s}, \mathbf{a})] - (\mathbb{E}_{\rho^{\pi}} [r(\mathbf{s}, \mathbf{a})] + \mathcal{H}(\pi)). \quad (\text{MaxEnt-IRL})$$

Diagram illustrating the components of the MaxEnt-IRL objective function:

- π_E : expert policy
- r : reward function
- ρ^{π} : occupancy measure (state-action visitation frequency)
- $\mathcal{H}(\pi)$: policy entropy

References:

[MaxEntIRL] Ziebart, B. D., Maas, A. L., Bagnell, J. A., & Dey, A. K. (2008, July). Maximum entropy inverse reinforcement learning. In AAAI (Vol. 8, pp. 1433-1438).

[ML-IRL] Zeng, S., Li, C., Garcia, A., & Hong, M. (2022). Maximum-likelihood inverse reinforcement learning with finite-time guarantees. Advances in Neural Information Processing Systems, 35, 10122-10135.

CONCEPTS OF INVERSE REINFORCEMENT LEARNING

ELEMENTS OF INVERSE REINFORCEMENT LEARNING

$$\min_{\pi \in \Pi} \max_{r \in \mathcal{R}} J(\pi_E, r) - J(\pi, r) = \mathbb{E}_{\rho^{\pi_E}}[r(\mathbf{s}, \mathbf{a})] - (\mathbb{E}_{\rho^{\pi}}[r(\mathbf{s}, \mathbf{a})] + \mathcal{H}(\pi)). \quad (\text{MaxEnt-IRL})$$

Diagram illustrating the components of the MaxEnt-IRL objective function:

- π_E : expert policy
- r : reward function
- ρ^{π} : occupancy measure (state-action visitation frequency)
- $\mathcal{H}(\pi)$: policy entropy

THE LIKELIHOOD OBJECTIVE OF INVERSE REINFORCEMENT LEARNING

$$\max_{\theta} \ell(\theta) := \mathbb{E}_{\rho^{\pi_E}}[\log \pi_{\theta}(\mathbf{a}|\mathbf{s})] \quad s.t. \quad \pi_{\theta} = \arg \max_{\pi} \mathbb{E}_{\pi}[r_{\theta}(\mathbf{s}, \mathbf{a})] \quad (\text{ML-IRL})$$

References:

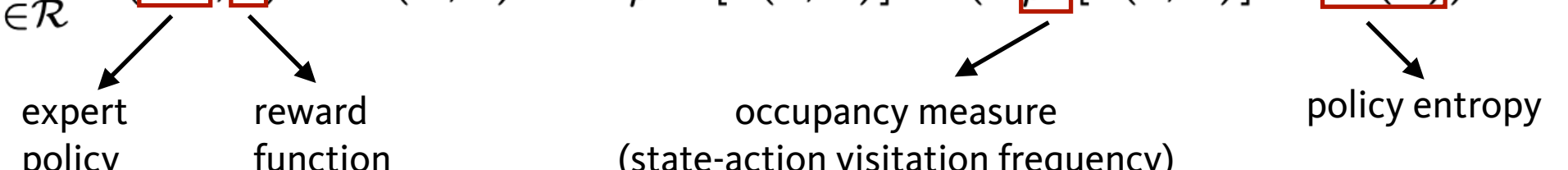
[MaxEntIRL] Ziebart, B. D., Maas, A. L., Bagnell, J. A., & Dey, A. K. (2008, July). Maximum entropy inverse reinforcement learning. In AAAI (Vol. 8, pp. 1433-1438).

[ML-IRL] Zeng, S., Li, C., Garcia, A., & Hong, M. (2022). Maximum-likelihood inverse reinforcement learning with finite-time guarantees. *Advances in Neural Information Processing Systems*, 35, 10122-10135.

CONCEPTS OF INVERSE REINFORCEMENT LEARNING

ELEMENTS OF INVERSE REINFORCEMENT LEARNING

$$\min_{\pi \in \Pi} \max_{r \in \mathcal{R}} J(\pi_E, r) - J(\pi, r) = \mathbb{E}_{\rho^{\pi_E}}[r(\mathbf{s}, \mathbf{a})] - (\mathbb{E}_{\rho^{\pi}}[r(\mathbf{s}, \mathbf{a})] + \mathcal{H}(\pi)). \quad (\text{MaxEnt-IRL})$$



THE LIKELIHOOD OBJECTIVE OF INVERSE REINFORCEMENT LEARNING

$$\max_{\theta} \ell(\theta) := \mathbb{E}_{\rho^{\pi_E}}[\log \pi_{\theta}(\mathbf{a}|\mathbf{s})] \quad s.t. \quad \pi_{\theta} = \arg \max_{\pi} \mathbb{E}_{\pi}[r_{\theta}(\mathbf{s}, \mathbf{a})] \quad (\text{ML-IRL})$$

THE EQUIVALENCE BETWEEN TWO OBJECTIVES

$$\begin{aligned} \ell(\theta) &= \mathbb{E}_{\rho^{\pi_E}}[r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\mathbf{s}_0 \sim \eta}[V_{r_{\theta}}^{\pi_{\theta}}(\mathbf{s}_0)] = J(\pi_E, r_{\theta}) - J(\pi_{\theta}, r_{\theta}), \\ \nabla_{\theta} \ell(\theta) &= \mathbb{E}_{\rho^{\pi_E}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\rho^{\pi_{\theta}}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})]. \end{aligned}$$

References:

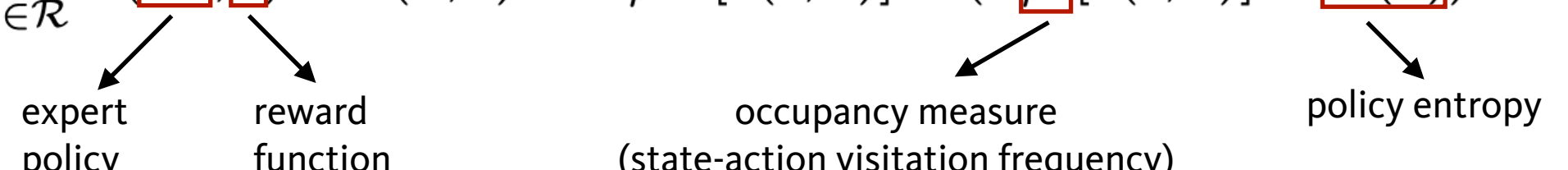
[MaxEntIRL] Ziebart, B. D., Maas, A. L., Bagnell, J. A., & Dey, A. K. (2008, July). Maximum entropy inverse reinforcement learning. In AAAI (Vol. 8, pp. 1433-1438).

[ML-IRL] Zeng, S., Li, C., Garcia, A., & Hong, M. (2022). Maximum-likelihood inverse reinforcement learning with finite-time guarantees. Advances in Neural Information Processing Systems, 35, 10122-10135.

CONCEPTS OF INVERSE REINFORCEMENT LEARNING

ELEMENTS OF INVERSE REINFORCEMENT LEARNING

$$\min_{\pi \in \Pi} \max_{r \in \mathcal{R}} J(\pi_E, r) - J(\pi, r) = \mathbb{E}_{\rho^{\pi_E}}[r(\mathbf{s}, \mathbf{a})] - (\mathbb{E}_{\rho^{\pi}}[r(\mathbf{s}, \mathbf{a})] + \mathcal{H}(\pi)). \quad (\text{MaxEnt-IRL})$$



THE LIKELIHOOD OBJECTIVE OF INVERSE REINFORCEMENT LEARNING

$$\max_{\theta} \ell(\theta) := \mathbb{E}_{\rho^{\pi_E}}[\log \pi_{\theta}(\mathbf{a}|\mathbf{s})] \quad s.t. \quad \pi_{\theta} = \arg \max_{\pi} \mathbb{E}_{\pi}[r_{\theta}(\mathbf{s}, \mathbf{a})] \quad (\text{ML-IRL})$$

THE EQUIVALENCE BETWEEN TWO OBJECTIVES

$$\begin{aligned} \ell(\theta) &= \mathbb{E}_{\rho^{\pi_E}}[r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\mathbf{s}_0 \sim \eta}[V_{r_{\theta}}^{\pi_{\theta}}(\mathbf{s}_0)] = J(\pi_E, r_{\theta}) - J(\pi_{\theta}, r_{\theta}), \\ \nabla_{\theta} \ell(\theta) &= \mathbb{E}_{\rho^{\pi_E}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\rho^{\pi_{\theta}}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})]. \end{aligned}$$

INSIGHT

The stability of IRL can be achieved by **monotonically** improving the **likelihood** of expert demonstrations, which is equivalent to reducing the **return gap** between the **expert policy** and the **reward-induced policy**.

References:

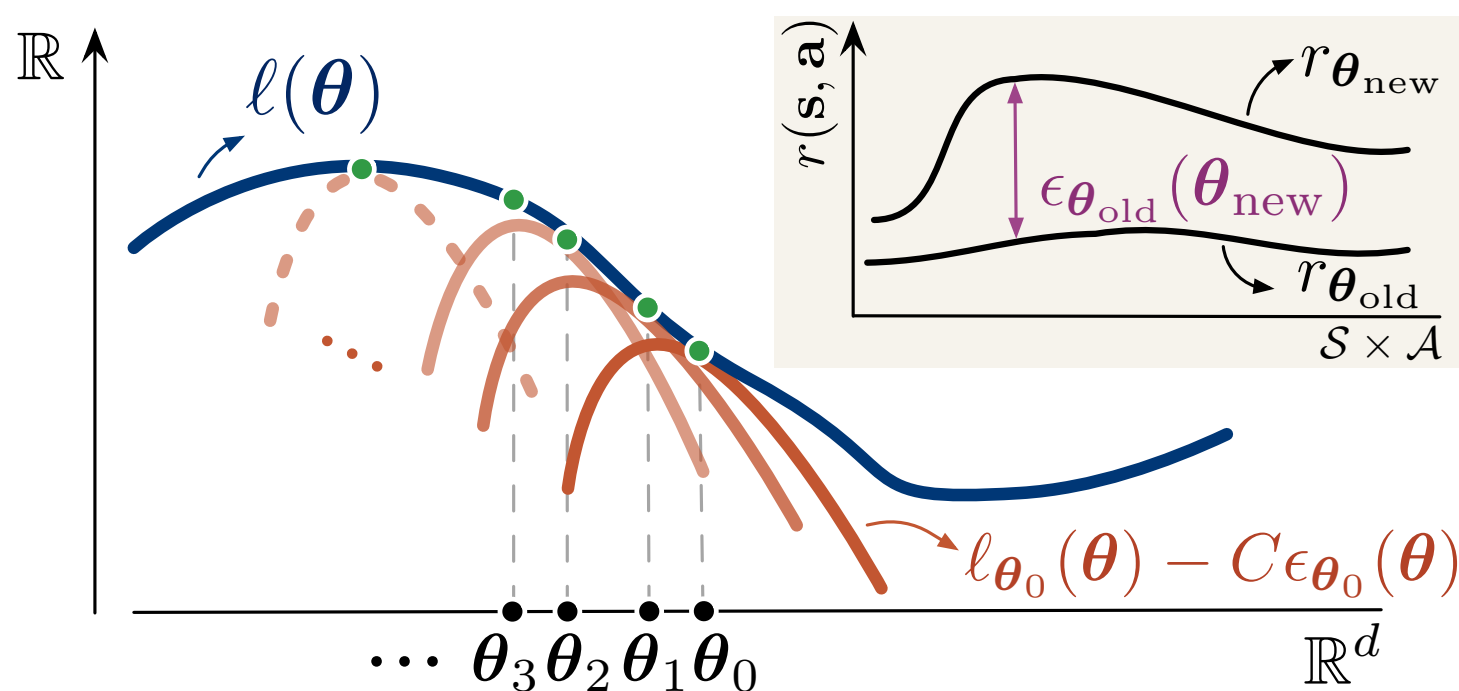
[MaxEntIRL] Ziebart, B. D., Maas, A. L., Bagnell, J. A., & Dey, A. K. (2008, July). Maximum entropy inverse reinforcement learning. In AAAI (Vol. 8, pp. 1433-1438).

[ML-IRL] Zeng, S., Li, C., Garcia, A., & Hong, M. (2022). Maximum-likelihood inverse reinforcement learning with finite-time guarantees. Advances in Neural Information Processing Systems, 35, 10122-10135.

TRUST REGION REWARD OPTIMIZATION (TRRO)

LOCAL APPROXIMATION OF THE LIKELIHOOD OBJECTIVE

$$\underbrace{\ell_{\theta_{\text{old}}}(\theta_{\text{old}}) = \ell(\theta_{\text{old}})}_{\equiv \mathbb{E}_{\rho^{\pi_E}}[r_{\theta_{\text{old}}}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\mathbf{s}_0 \sim \eta}[V_{r_{\theta_{\text{old}}}}^{\pi_{\text{old}}}(\mathbf{s}_0)]} \quad \text{and} \quad \underbrace{\nabla_{\theta} \ell_{\theta_{\text{old}}}(\theta)|_{\theta=\theta_{\text{old}}} = \nabla_{\theta} \ell(\theta)|_{\theta=\theta_{\text{old}}}}_{\equiv \mathbb{E}_{\rho^{\pi_E}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\rho^{\pi_{\text{old}}}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})]|_{\theta=\theta_{\text{old}}}} .$$



TRUST REGION REWARD OPTIMIZATION (TRRO)

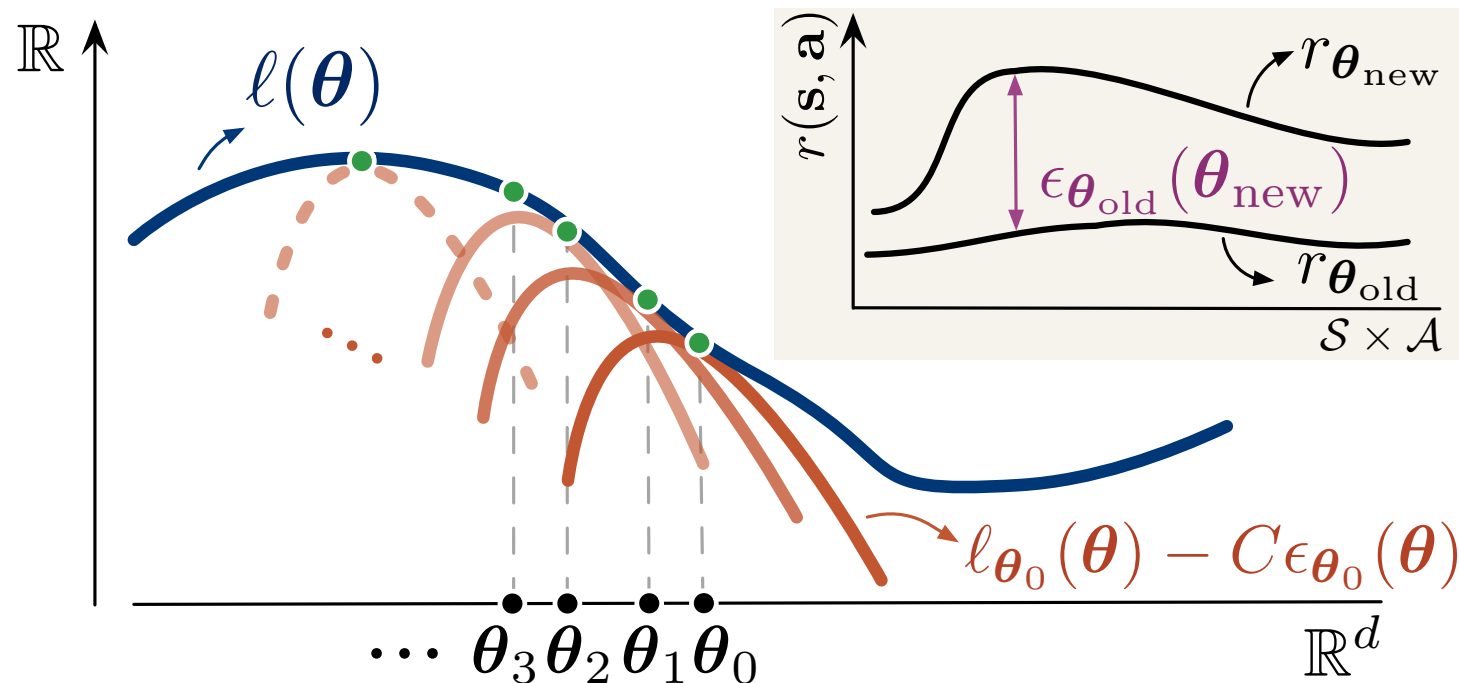
LOCAL APPROXIMATION OF THE LIKELIHOOD OBJECTIVE

$$\underbrace{\ell_{\theta_{\text{old}}}(\theta_{\text{old}}) = \ell(\theta_{\text{old}})}_{\equiv \mathbb{E}_{\rho^{\pi_E}}[r_{\theta_{\text{old}}}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\mathbf{s}_0 \sim \eta}[V_{r_{\theta_{\text{old}}}}^{\pi_{\text{old}}}(\mathbf{s}_0)]} \quad \text{and} \quad \underbrace{\nabla_{\theta} \ell_{\theta_{\text{old}}}(\theta)|_{\theta=\theta_{\text{old}}} = \nabla_{\theta} \ell(\theta)|_{\theta=\theta_{\text{old}}}}_{\equiv \mathbb{E}_{\rho^{\pi_E}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\rho^{\pi_{\text{old}}}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})]|_{\theta=\theta_{\text{old}}}} .$$

THE GUARANTEE OF MONOTONIC IMPROVEMENT OF THE LIKELIHOOD

$\ell(\theta_{\text{new}}) \geq \ell_{\theta_{\text{old}}}(\theta_{\text{new}}) - C\epsilon_{\theta_{\text{old}}}(\theta_{\text{new}})$, where the constant

$$C = \frac{2|\mathcal{A}|}{(1-\gamma)^2} + \frac{(5-\gamma)|\mathcal{A}|R + (\gamma - \gamma^2 + 2)|\mathcal{A}| \log |\mathcal{A}|}{(1-\gamma)^4}.$$



TRUST REGION REWARD OPTIMIZATION (TRRO)

LOCAL APPROXIMATION OF THE LIKELIHOOD OBJECTIVE

$$\underbrace{\ell_{\theta_{\text{old}}}(\theta_{\text{old}}) = \ell(\theta_{\text{old}})}_{\equiv \mathbb{E}_{\rho^{\pi_E}}[r_{\theta_{\text{old}}}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\mathbf{s}_0 \sim \eta}[V_{r_{\theta_{\text{old}}}^{\pi_{\text{old}}}}(\mathbf{s}_0)]} \quad \text{and} \quad \underbrace{\nabla_{\theta} \ell_{\theta_{\text{old}}}(\theta)|_{\theta=\theta_{\text{old}}} = \nabla_{\theta} \ell(\theta)|_{\theta=\theta_{\text{old}}}}_{\equiv \mathbb{E}_{\rho^{\pi_E}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})] - \mathbb{E}_{\rho^{\pi_{\text{old}}}}[\nabla_{\theta} r_{\theta}(\mathbf{s}, \mathbf{a})]|_{\theta=\theta_{\text{old}}}}.$$

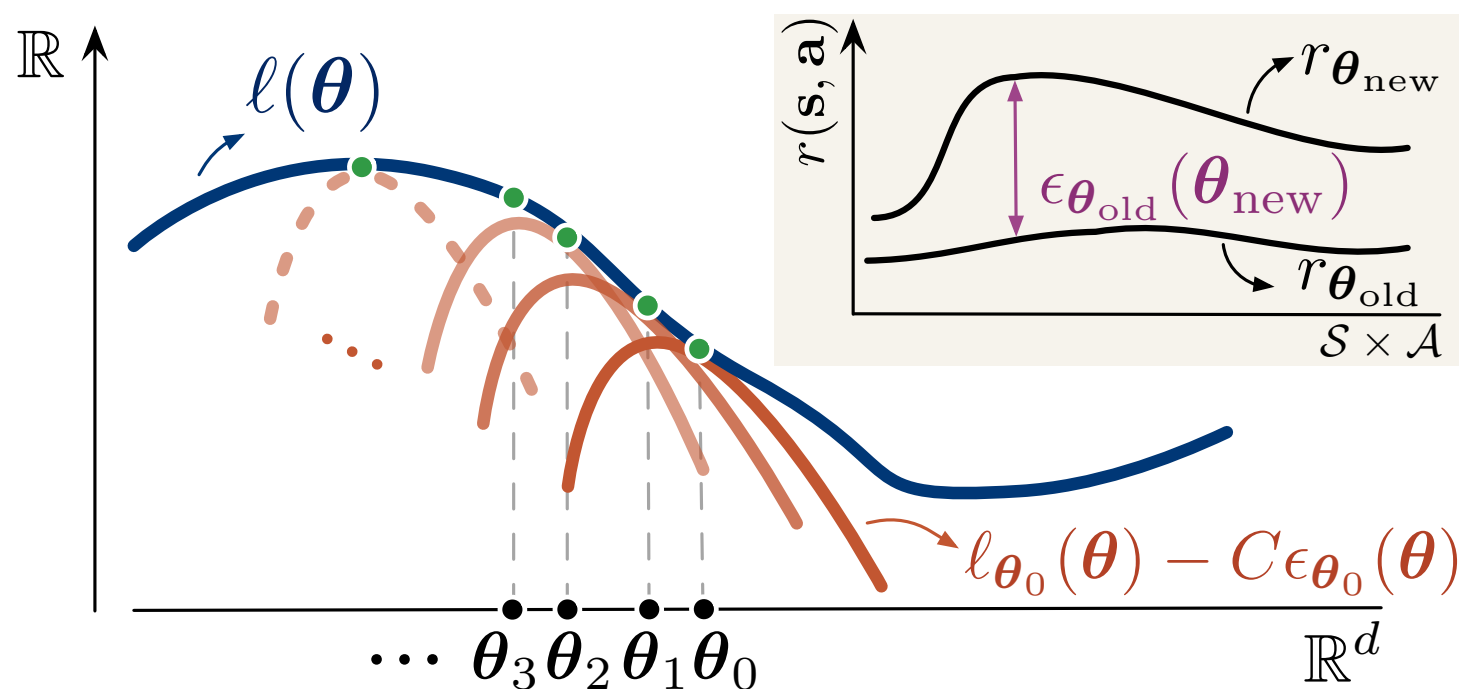
THE GUARANTEE OF MONOTONIC IMPROVEMENT OF THE LIKELIHOOD

$\ell(\theta_{\text{new}}) \geq \ell_{\theta_{\text{old}}}(\theta_{\text{new}}) - C\epsilon_{\theta_{\text{old}}}(\theta_{\text{new}})$, where the constant

$$C = \frac{2|\mathcal{A}|}{(1-\gamma)^2} + \frac{(5-\gamma)|\mathcal{A}|R + (\gamma - \gamma^2 + 2)|\mathcal{A}| \log |\mathcal{A}|}{(1-\gamma)^4}.$$

THE PROCEDURE OF TRRO

$$\pi = \arg \max_{\pi} J(\pi, r_{\theta_{\text{old}}}) \xrightarrow[\theta_{\text{old}} \leftarrow \theta_{\text{new}}]{\pi \rightarrow \pi_{\text{old}}} \theta_{\text{new}} = \arg \max_{\theta} \ell_{\theta_{\text{old}}}(\theta) - C\epsilon_{\theta_{\text{old}}}(\theta). \quad (\text{TRRO})$$



PROXIMAL INVERSE REWARD OPTIMIZATION ALGORITHM (PIRO)

ADAPTIVE REWARD UPDATE

$$\bar{\epsilon}_{\theta_{\text{old}}}(\boldsymbol{\theta}) := \left(\sum_{(\mathbf{s}, \mathbf{a}) \in \hat{\mathcal{D}}_E \cup \mathcal{D}_S} (r_{\theta_{\text{old}}}(\mathbf{s}, \mathbf{a}) - r_{\boldsymbol{\theta}}(\mathbf{s}, \mathbf{a}))^2 \right)^{1/2}$$

$$\boldsymbol{\theta}_{\text{new}} = \arg \max_{\boldsymbol{\theta}} L_{\theta_{\text{old}}}(\boldsymbol{\theta}) := \ell_{\theta_{\text{old}}}(\boldsymbol{\theta}) - \mu \bar{\epsilon}_{\theta_{\text{old}}}(\boldsymbol{\theta}).$$

If $\bar{\epsilon}_{\theta_{\text{old}}}(\boldsymbol{\theta}) > \bar{\epsilon}^{\text{target}} \times x$, then $\mu \leftarrow \mu \times y$; If $\bar{\epsilon}_{\theta_{\text{old}}}(\boldsymbol{\theta}) < \bar{\epsilon}^{\text{target}}/x$, then $\mu \leftarrow \mu/y$,

PROXIMAL INVERSE REWARD OPTIMIZATION ALGORITHM (PIRO)

ADAPTIVE REWARD UPDATE

$$\bar{\epsilon}_{\theta_{\text{old}}}(\theta) := \left(\sum_{(\mathbf{s}, \mathbf{a}) \in \hat{\mathcal{D}}_E \cup \mathcal{D}_S} (r_{\theta_{\text{old}}}(\mathbf{s}, \mathbf{a}) - r_{\theta}(\mathbf{s}, \mathbf{a}))^2 \right)^{1/2}$$

$$\theta_{\text{new}} = \arg \max_{\theta} L_{\theta_{\text{old}}}(\theta) := \ell_{\theta_{\text{old}}}(\theta) - \mu \bar{\epsilon}_{\theta_{\text{old}}}(\theta).$$

If $\bar{\epsilon}_{\theta_{\text{old}}}(\theta) > \bar{\epsilon}^{\text{target}} \times x$, then $\mu \leftarrow \mu \times y$; If $\bar{\epsilon}_{\theta_{\text{old}}}(\theta) < \bar{\epsilon}^{\text{target}}/x$, then $\mu \leftarrow \mu/y$,

PRACTICAL POLICY OPTIMIZATION

$\pi \leftarrow k$ SAC rounds with $r_{\theta_{\text{old}}}, \pi_{\text{old}}$

PROXIMAL INVERSE REWARD OPTIMIZATION ALGORITHM (PIRO)

ADAPTIVE REWARD UPDATE

$$\bar{\epsilon}_{\theta_{\text{old}}}(\theta) := \left(\sum_{(\mathbf{s}, \mathbf{a}) \in \hat{\mathcal{D}}_E \cup \mathcal{D}_S} (r_{\theta_{\text{old}}}(\mathbf{s}, \mathbf{a}) - r_{\theta}(\mathbf{s}, \mathbf{a}))^2 \right)^{1/2}$$

$$\theta_{\text{new}} = \arg \max_{\theta} L_{\theta_{\text{old}}}(\theta) := \ell_{\theta_{\text{old}}}(\theta) - \mu \bar{\epsilon}_{\theta_{\text{old}}}(\theta).$$

If $\bar{\epsilon}_{\theta_{\text{old}}}(\theta) > \bar{\epsilon}^{\text{target}} \times x$, then $\mu \leftarrow \mu \times y$; If $\bar{\epsilon}_{\theta_{\text{old}}}(\theta) < \bar{\epsilon}^{\text{target}}/x$, then $\mu \leftarrow \mu/y$,

PRACTICAL POLICY OPTIMIZATION

$\pi \leftarrow k$ SAC rounds with $r_{\theta_{\text{old}}}, \pi_{\text{old}}$

PIRO ALGORITHM

$\pi \leftarrow k$ SAC rounds with $r_{\theta_{\text{old}}}, \pi_{\text{old}}$ $\xrightleftharpoons[\theta_{\text{old}} \leftarrow \theta_{\text{new}}]{\pi \rightarrow \pi_{\text{old}}}$ $\theta_{\text{new}} \leftarrow n$ grad. steps with $\nabla_{\theta} L_{\theta_{\text{old}}}(\theta)$. (PIRO)

PROXIMAL INVERSE REWARD OPTIMIZATION ALGORITHM (PIRO)

ADAPTIVE REWARD UPDATE

$$\bar{\epsilon}_{\theta_{\text{old}}}(\theta) := \left(\sum_{(s,a) \in \hat{\mathcal{D}}_E \cup \mathcal{D}_S} (r_{\theta_{\text{old}}}(s,a) - r_{\theta}(s,a))^2 \right)^{1/2}$$

$$\theta_{\text{new}} = \arg \max_{\theta} L_{\theta_{\text{old}}}(\theta) := \ell_{\theta_{\text{old}}}(\theta) - \mu \bar{\epsilon}_{\theta_{\text{old}}}(\theta).$$

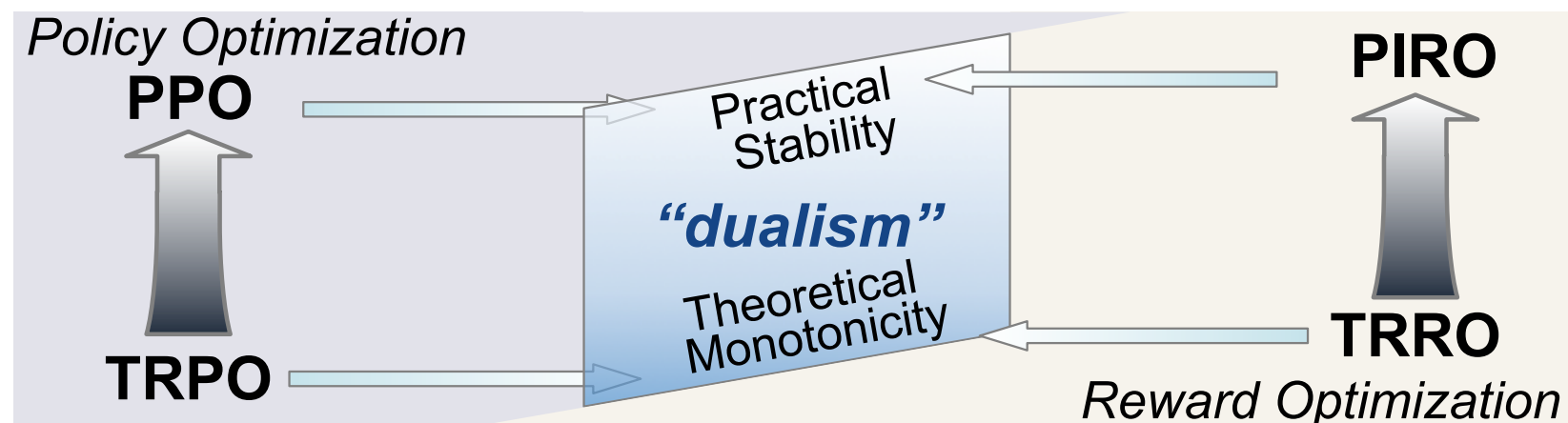
If $\bar{\epsilon}_{\theta_{\text{old}}}(\theta) > \bar{\epsilon}^{\text{target}} \times x$, then $\mu \leftarrow \mu \times y$; If $\bar{\epsilon}_{\theta_{\text{old}}}(\theta) < \bar{\epsilon}^{\text{target}}/x$, then $\mu \leftarrow \mu/y$,

PRACTICAL POLICY OPTIMIZATION

$\pi \leftarrow k$ SAC rounds with $r_{\theta_{\text{old}}}, \pi_{\text{old}}$

PIRO ALGORITHM

$\pi \leftarrow k$ SAC rounds with $r_{\theta_{\text{old}}}, \pi_{\text{old}}$ $\xrightleftharpoons[\theta_{\text{old}} \leftarrow \theta_{\text{new}}]{\pi \rightarrow \pi_{\text{old}}}$ $\theta_{\text{new}} \leftarrow n$ grad. steps with $\nabla_{\theta} L_{\theta_{\text{old}}}(\theta)$. (PIRO)



REWARD RECOVERY AND POLICY IMITATION

Table 1: Averaged Rewards (five independent runs) on five MuJoCo and four Gym Robotics tasks.

	Task	Expert	IL		Adv. IRL (Online)				Adv. (Offline)	Non-Adv. Online			Non-Adv. Offline		PIRO	Gain
			BC	GAIL	MM	AIRL	FILTER	HyPE	DAC	IQ	ML-IRL	f -IRL	CSIL	P ² IL		
MuJoCo	Ant-v4	5926.2	1631.5	996.9	-304.0	991.4	-376.3	2800.5	923.8	3589.8	5382.5	980.4	420.7	976.6	5967.2	+584.7
	Humanoid-v4	5501.0	418.1	508.4	367.2	281.4	291.7	717.5	76.3	1847.5	5573.4	470.4	–	–	5954.9	+381.5
	Walker2d-v4	5524.5	384.4	4158.1	70.4	72.8	77.7	1478.7	-3.0	3023.0	4794.7	243.8	686.1	1054.0	5643.7	+849.0
	Hopper-v4	3632.8	1034.4	3535.7	57.8	13.5	37.3	2593.7	3321.6	3424.5	3316.4	361.7	6.7	25.8	3362.0	-173.7
	Halfcheetah-v4	12266.1	221.2	1298.8	20.3	2251.4	0.3	6473.4	9645.0	3825.5	11873.2	-0.7	-107.2	-0.1	12587.4	+714.2
Robotics	AntMaze-UMazeDense-v4	35.6	8.8	5.2	5.1	4.5	6.1	11.9	–	3.9	4.2	3.6	–	3.4	25.7	+13.8
	AntMaze-MediumDense-v4	26.9	1.1	1.3	3.4	2.6	1.9	3.0	–	3.4	0.9	1.1	–	2.9	9.4	+6.0
	AntMaze-LargeDense-v4	11.5	1.1	0.9	1.7	3.4	0.6	1.5	–	0.8	0.3	0.9	–	0.2	8.8	+5.4
	AdroitHandePen-Human-v1	1062.5	44.1	-8.7	-344.3	-593.9	-685.4	-866.7	–	-751.9	-251.2	-65.3	–	-61.2	254.0	+209.9
	runtime per iteration	-	-	3-14s	8-79s	5-8s	9-41s	11-70s	135-142s	7-57s	93-166s	16-85s	68-90s	20-111s	96-178s	–

Note: DAC, CSIL and P²IL are not evaluated on certain tasks due to compatibility issues caused by version conflicts. Specifically, the current implementations of DAC and P²IL are incompatible with the current Gymnasium Robotics suite, while P²IL and CSIL are incompatible with the Humanoid version used in testing other algorithms.

EXPERIMENTAL EVALUATION

REWARD RECOVERY AND POLICY IMITATION

Table 1: Averaged Rewards (five independent runs) on five MuJoCo and four Gym Robotics tasks.

	Task	Expert	IL		Adv. IRL (Online)				Adv. (Offline)	Non-Adv. Online			Non-Adv. Offline		PIRO	Gain
			BC	GAIL	MM	AIRL	FILTER	HyPE	DAC	IQ	ML-IRL	f -IRL	CSIL	P ² IL		
MuJoCo	Ant-v4	5926.2	1631.5	996.9	-304.0	991.4	-376.3	2800.5	923.8	3589.8	5382.5	980.4	420.7	976.6	5967.2	+584.7
	Humanoid-v4	5501.0	418.1	508.4	367.2	281.4	291.7	717.5	76.3	1847.5	5573.4	470.4	–	–	5954.9	+381.5
	Walker2d-v4	5524.5	384.4	4158.1	70.4	72.8	77.7	1478.7	-3.0	3023.0	4794.7	243.8	686.1	1054.0	5643.7	+849.0
	Hopper-v4	3632.8	1034.4	3535.7	57.8	13.5	37.3	2593.7	3321.6	3424.5	3316.4	361.7	6.7	25.8	3362.0	-173.7
	Halfcheetah-v4	12266.1	221.2	1298.8	20.3	2251.4	0.3	6473.4	9645.0	3825.5	11873.2	-0.7	-107.2	-0.1	12587.4	+714.2
Robotics	AntMaze-UMazeDense-v4	35.6	8.8	5.2	5.1	4.5	6.1	11.9	–	3.9	4.2	3.6	–	3.4	25.7	+13.8
	AntMaze-MediumDense-v4	26.9	1.1	1.3	3.4	2.6	1.9	3.0	–	3.4	0.9	1.1	–	2.9	9.4	+6.0
	AntMaze-LargeDense-v4	11.5	1.1	0.9	1.7	3.4	0.6	1.5	–	0.8	0.3	0.9	–	0.2	8.8	+5.4
	AdroitHandPen-Human-v1	1062.5	44.1	-8.7	-344.3	-593.9	-685.4	-866.7	–	-751.9	-251.2	-65.3	–	-61.2	254.0	+209.9
	runtime per iteration	–	–	3-14s	8-79s	5-8s	9-41s	11-70s	135-142s	7-57s	93-166s	16-85s	68-90s	20-111s	96-178s	–

Note: DAC, CSIL and P²IL are not evaluated on certain tasks due to compatibility issues cause by version conflicts. Specifically, the current implementations of DAC and P²IL are incompatible with the current Gymnasium Robotics suite, while P²IL and CSIL are incompatible with the Humanoid version used in testing other algorithms.

LEARNING STABILITY

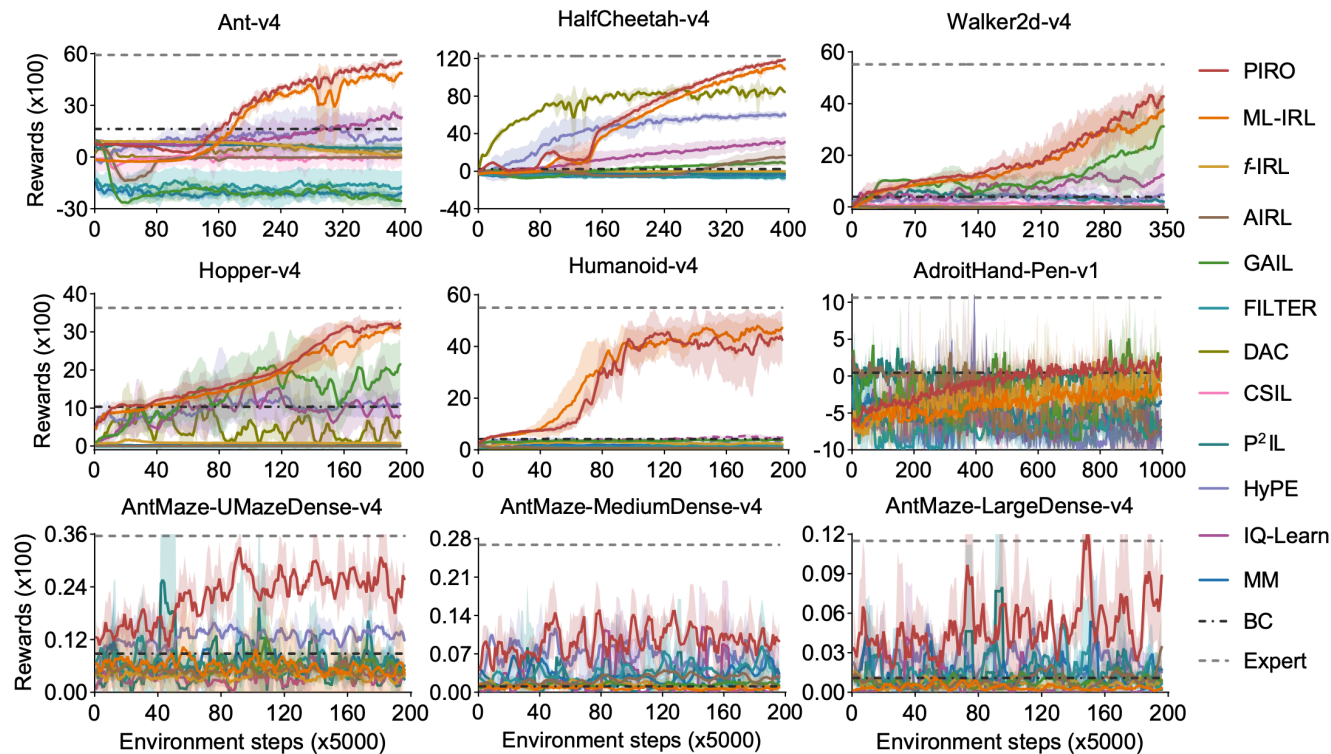


Figure 4: Reward curves of algorithms on MuJoCo locomotion tasks and Gym Robotics tasks.

EXPERIMENTAL EVALUATION

REWARD RECOVERY AND POLICY IMITATION

Table 1: Averaged Rewards (five independent runs) on five MuJoCo and four Gym Robotics tasks.

	Task	Expert	IL		Adv. IRL (Online)				Adv. (Offline)	Non-Adv. Online			Non-Adv. Offline		PIRO	Gain
			BC	GAIL	MM	AIRL	FILTER	HyPE	DAC	IQ	ML-IRL	f -IRL	CSIL	P ² IL		
MuJoCo	Ant-v4	5926.2	1631.5	996.9	-304.0	991.4	-376.3	2800.5	923.8	3589.8	5382.5	980.4	420.7	976.6	5967.2	+584.7
	Humanoid-v4	5501.0	418.1	508.4	367.2	281.4	291.7	717.5	76.3	1847.5	5573.4	470.4	—	—	5954.9	+381.5
	Walker2d-v4	5524.5	384.4	4158.1	70.4	72.8	77.7	1478.7	-3.0	3023.0	4794.7	243.8	686.1	1054.0	5643.7	+849.0
	Hopper-v4	3632.8	1034.4	3535.7	57.8	13.5	37.3	2593.7	3321.6	3424.5	3316.4	361.7	6.7	25.8	3362.0	-173.7
	Halfcheetah-v4	12266.1	221.2	1298.8	20.3	2251.4	0.3	6473.4	9645.0	3825.5	11873.2	-0.7	-107.2	-0.1	12587.4	+714.2
Robotics	AntMaze-UMazeDense-v4	35.6	8.8	5.2	5.1	4.5	6.1	11.9	—	3.9	4.2	3.6	—	3.4	25.7	+13.8
	AntMaze-MediumDense-v4	26.9	1.1	1.3	3.4	2.6	1.9	3.0	—	3.4	0.9	1.1	—	2.9	9.4	+6.0
	AntMaze-LargeDense-v4	11.5	1.1	0.9	1.7	3.4	0.6	1.5	—	0.8	0.3	0.9	—	0.2	8.8	+5.4
	AdroitHandPen-Human-v1	1062.5	44.1	-8.7	-344.3	-593.9	-685.4	-866.7	—	-751.9	-251.2	-65.3	—	-61.2	254.0	+209.9
	runtime per iteration	—	—	3-14s	8-79s	5-8s	9-41s	11-70s	135-142s	7-57s	93-166s	16-85s	68-90s	20-111s	96-178s	—

Note: DAC, CSIL and P²IL are not evaluated on certain tasks due to compatibility issues caused by version conflicts. Specifically, the current implementations of DAC and P²IL are incompatible with the current Gymnasium Robotics suite, while P²IL and CSIL are incompatible with the Humanoid version used in testing other algorithms.

STATE-ONLY REWARDS

LEARNING STABILITY

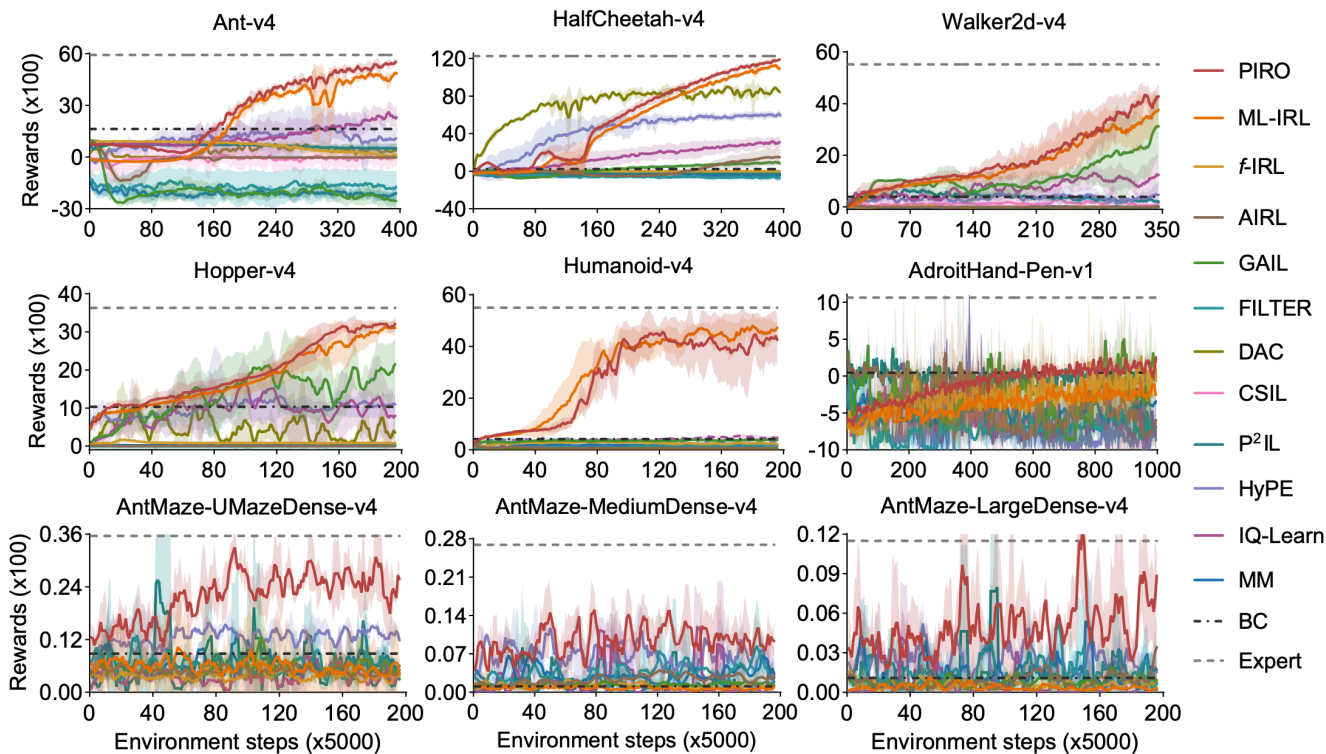


Figure 4: Reward curves of algorithms on MuJoCo locomotion tasks and Gym Robotics tasks.

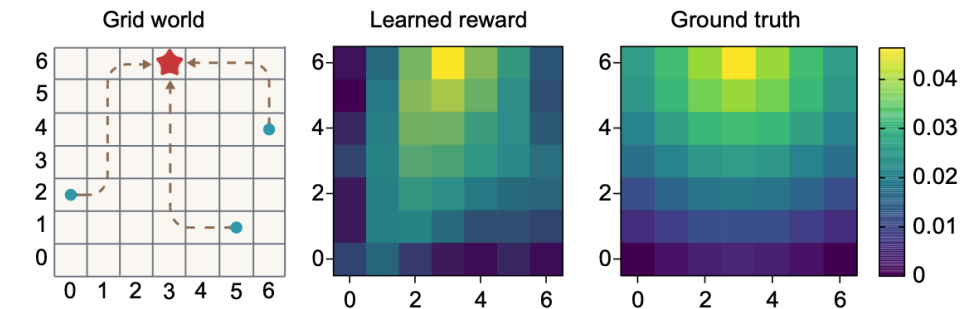


Figure 5: Experiments on reward recovery in tasks with state-only rewards. **Left:** The task is a 7×7 grid world, where the agent starts from a random initial position (blue circles) with the objective of reaching the target position (red star) via the shortest possible path. **Right:** The ground truth reward at each position is defined as the negative Euclidean distance to the terminal state. **Middle:** The reward recovered by PIRO and the ground-truth reward function is highly consistent with the ground truth reward. Cumulative rewards: -9.24 (expert) vs. -8.48 (PIRO).

REWARD TRANSFER

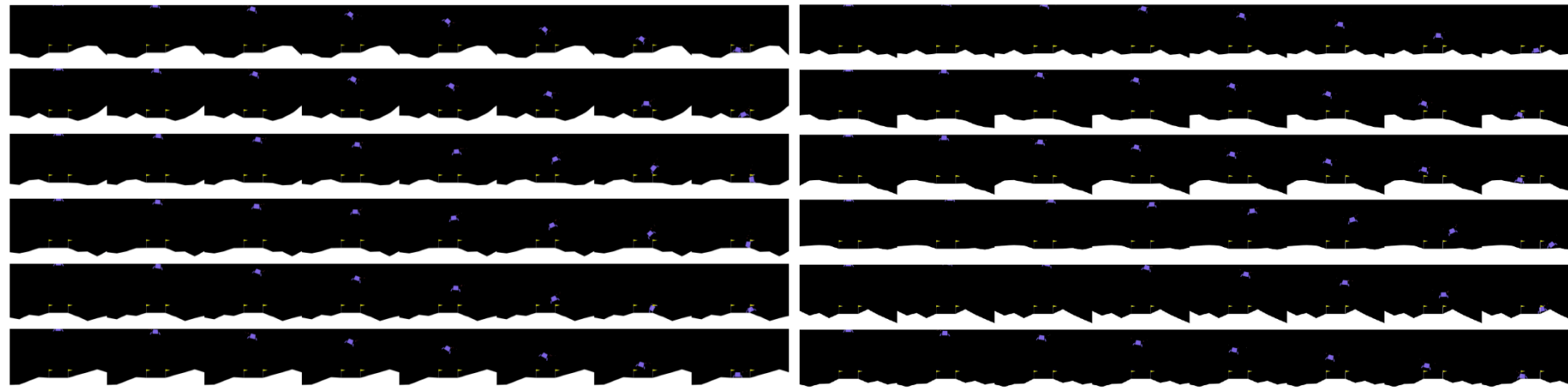


Figure 6: **Results for reward transfer to new environments with altered dynamics.** **Left panels:** Policy behavior learned by PIRO in the original LunarLander environment. PIRO succeeds in most cases. **Right panels:** Policy behavior under PIRO's learned reward function in LunarLander with altered dynamics (stochastic wind added). The policy is robust in general, despite some failure cases, e.g., row 3.

REWARD TRANSFER

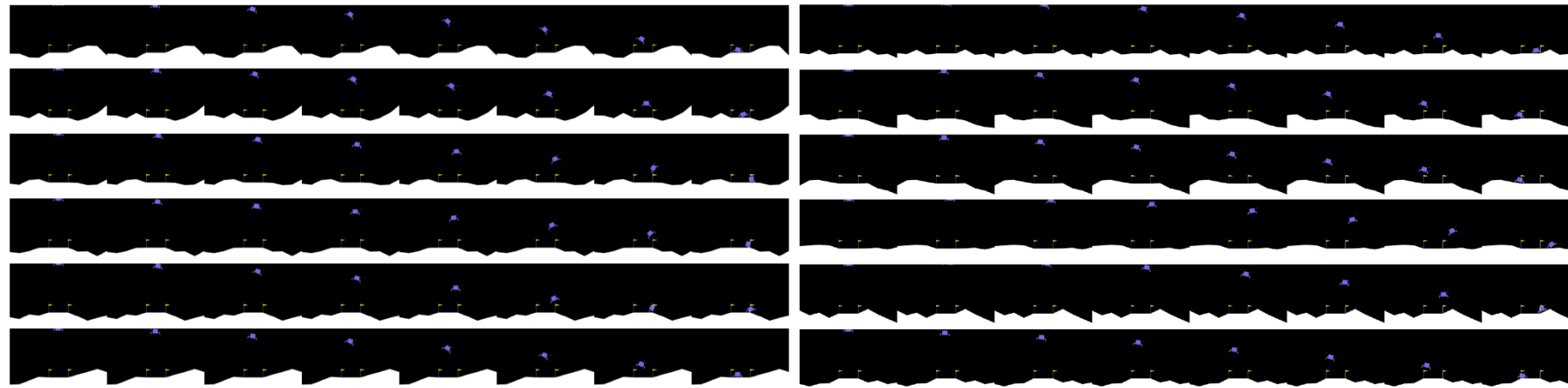


Figure 6: **Results for reward transfer to new environments with altered dynamics.** **Left panels:** Policy behavior learned by PIRO in the original LunarLander environment. PIRO succeeds in most cases. **Right panels:** Policy behavior under PIRO's learned reward function in LunarLander with altered dynamics (stochastic wind added). The policy is robust in general, despite some failure cases, e.g., row 3.

REAL-WORLD CASE STUDY— ANIMAL BEHAVIOR MODELING



Figure 10: Fifteen types of the meerkat behaviors.

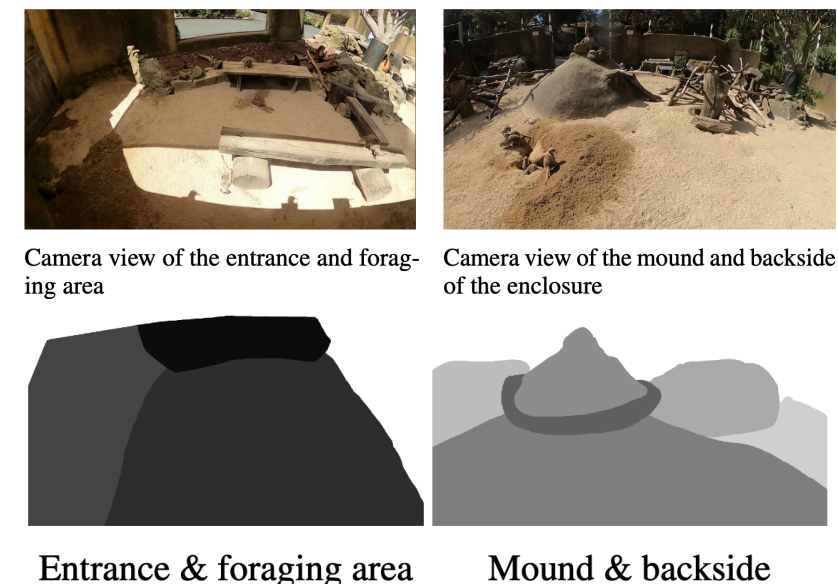


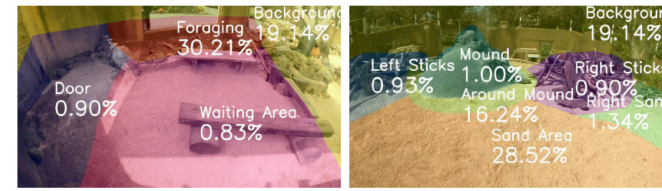
Figure 12: Different colors are labelled for each area to visually illustrate the division of meerkat activity zones.

EXPERIMENTAL EVALUATION (CON'D)

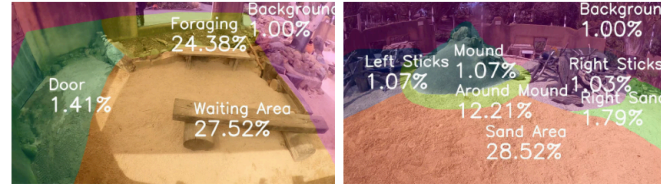
REAL-WORLD CASE STUDY— ANIMAL BEHAVIOR MODELING



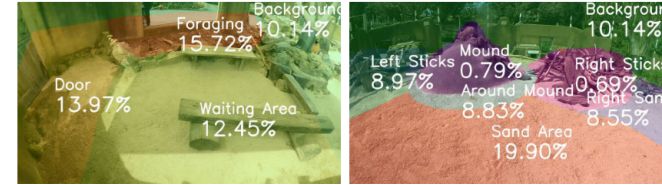
(a) Expert



(b) PIRO (weighted mean error: 5.5%)



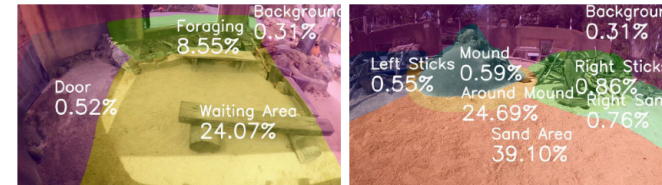
(c) AIRL (weighted mean error: 8.7%)



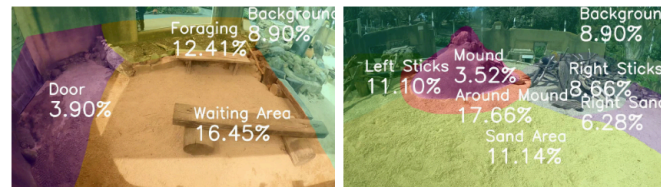
(d) BC (weighted mean error: 9.3%)



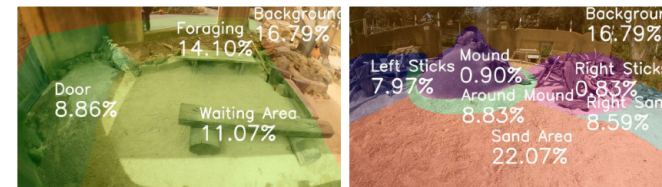
(e) GAIL (weighted mean error: 11.8%)



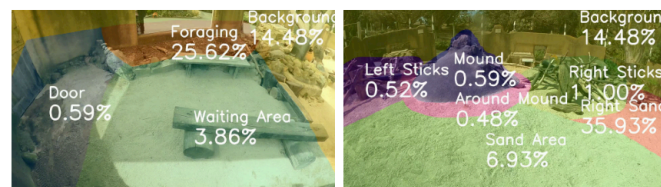
(f) f -IRL (weighted mean error: 16.1%)



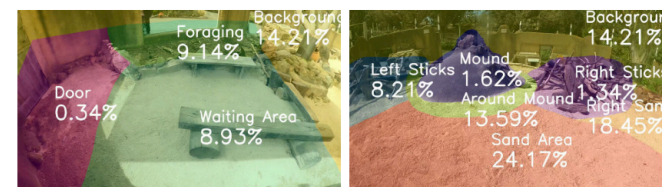
(g) FILTER (weighted mean error: 10.0%)



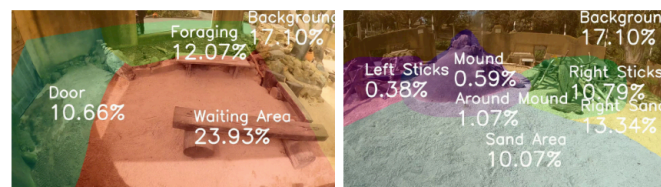
(h) MM (weighted mean error: 10.5%)



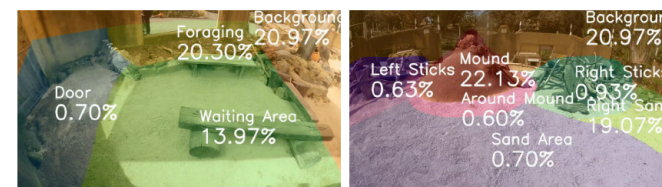
(i) IQ-Learn (weighted mean error: 9.8%)



(j) ML-IRL (weighted mean error: 11.3%)

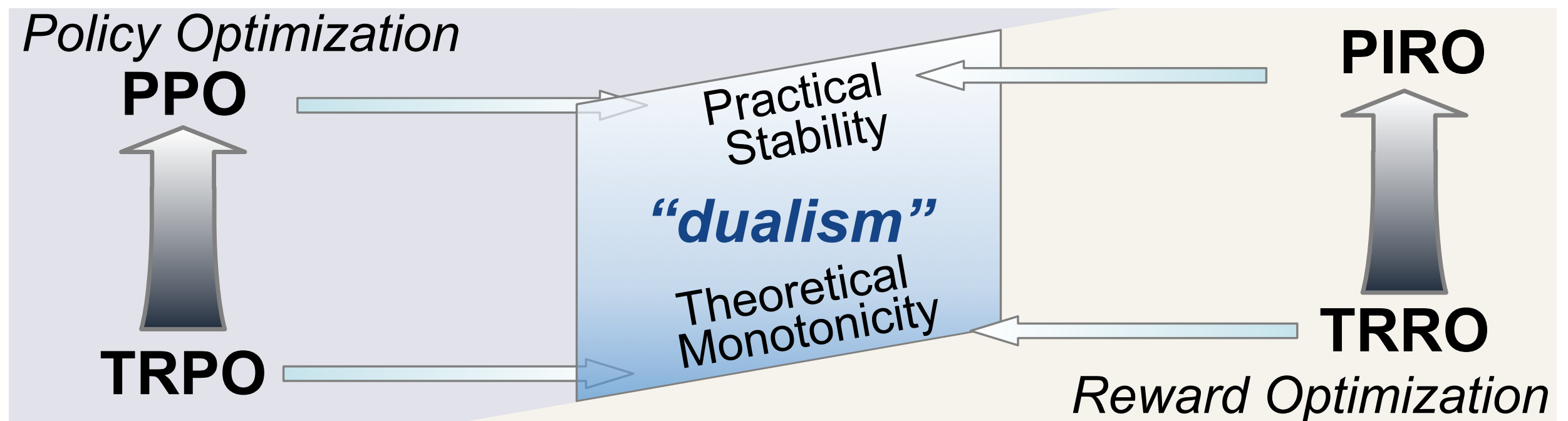


(k) HyPE (weighted mean error: 14.6%)



(l) P²IL (weighted mean error: 13.0%)

Figure 13: **Regional visitation frequency map** generated by analyzing real meerkat trajectories alongside those produced by algorithms. PIRO achieves the lowest weighted mean error.



Thanks for your watching !

Scan to View the Project

