

Zeroth-Order Policy Gradient for Reinforcement Learning from Human Feedback without Reward Inference

Qining Zhang

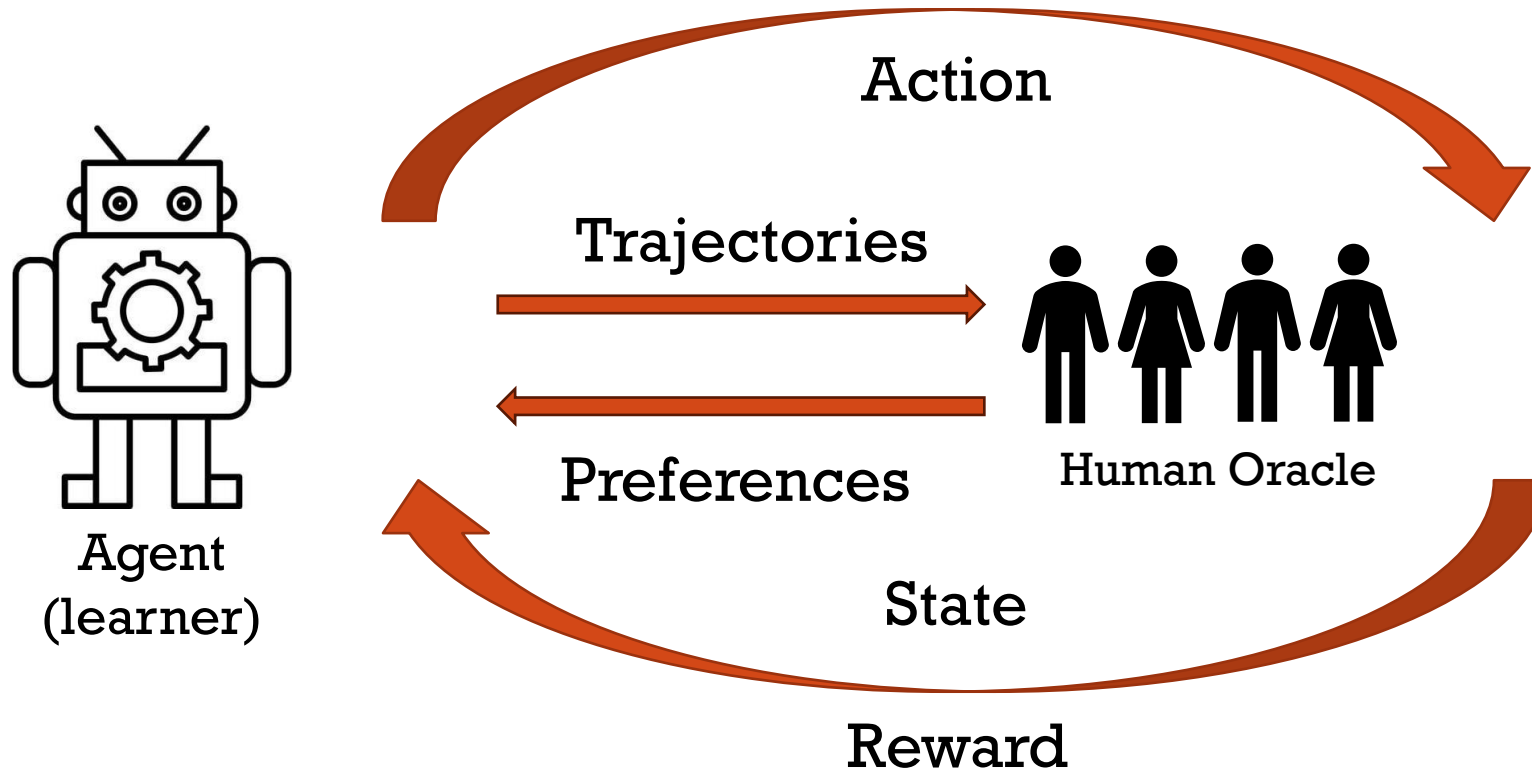
EECS, University of Michigan, Ann Arbor

Joint work with Lei Ying(UMich)

RL from Human Feedback

- Learn from Reward v.s. Learn from Preference

- Inverse RL (Russell 1998, Ng&Russell 2000, ...)
- Dueling bandits (Yue&Joachims 2000, Ailon, Karnin&Joachims, 2014, ...)
- Preference-based RL (Akrour et al. 2011, Cheng et al. 2011, ...)



Environment

Applications

Large Language Models



Preferences of Prompt Responses

Video Game Bot



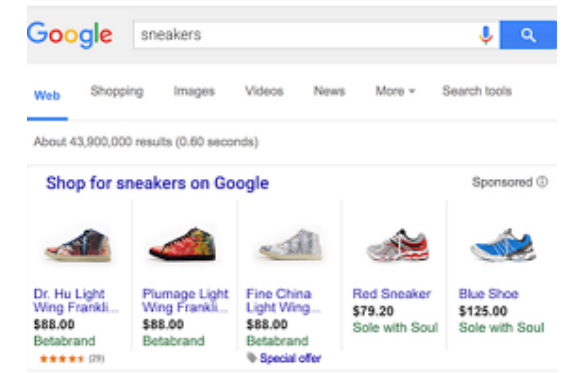
Preferences of Gaming Strategies

Robotics



Preferences of Motions

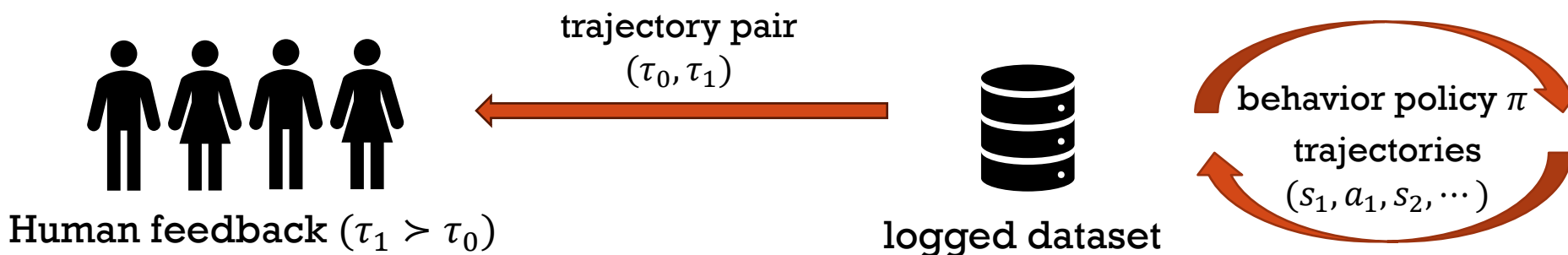
Recommendation Systems



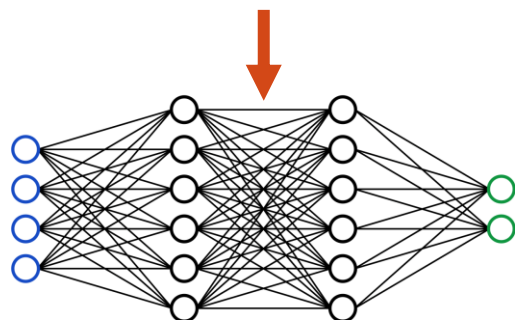
Preferences of Merchandise

RLHF: The GPT-4 Approach

- Bradley-Terry (BT) Preference Model + Reward Inference + PPO (Achiam et al. 2023)

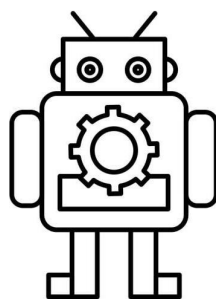


$$\text{BT: } \mathbb{P}(\tau_1 > \tau_0) = (1 + e^{r(\tau_0) - r(\tau_1)})^{-1}$$



Reward model: $r_h(s, a)$

Reward inference



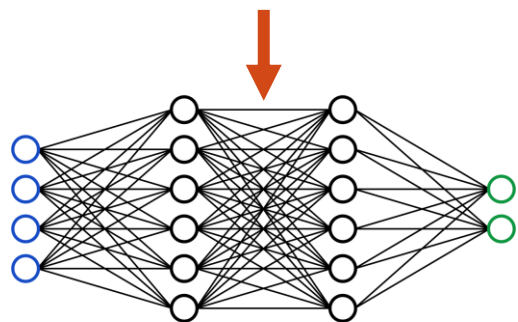
Proximal Policy Optimization (PPO)
(Schulman et al. 2017)

RLHF: The GPT-4 Approach

- Bradley-Terry (BT) Preference Model + Reward Inference + PPO (Achiam et al. 2023)

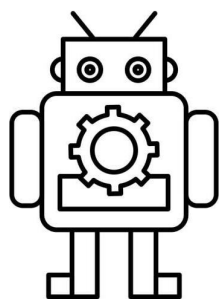


$$\text{BT: } \mathbb{P}(\tau_1 > \tau_0) = (1 + e^{r(\tau_0) - r(\tau_1)})^{-1}$$



Reward model: $r_h(s, a)$

Reward inference



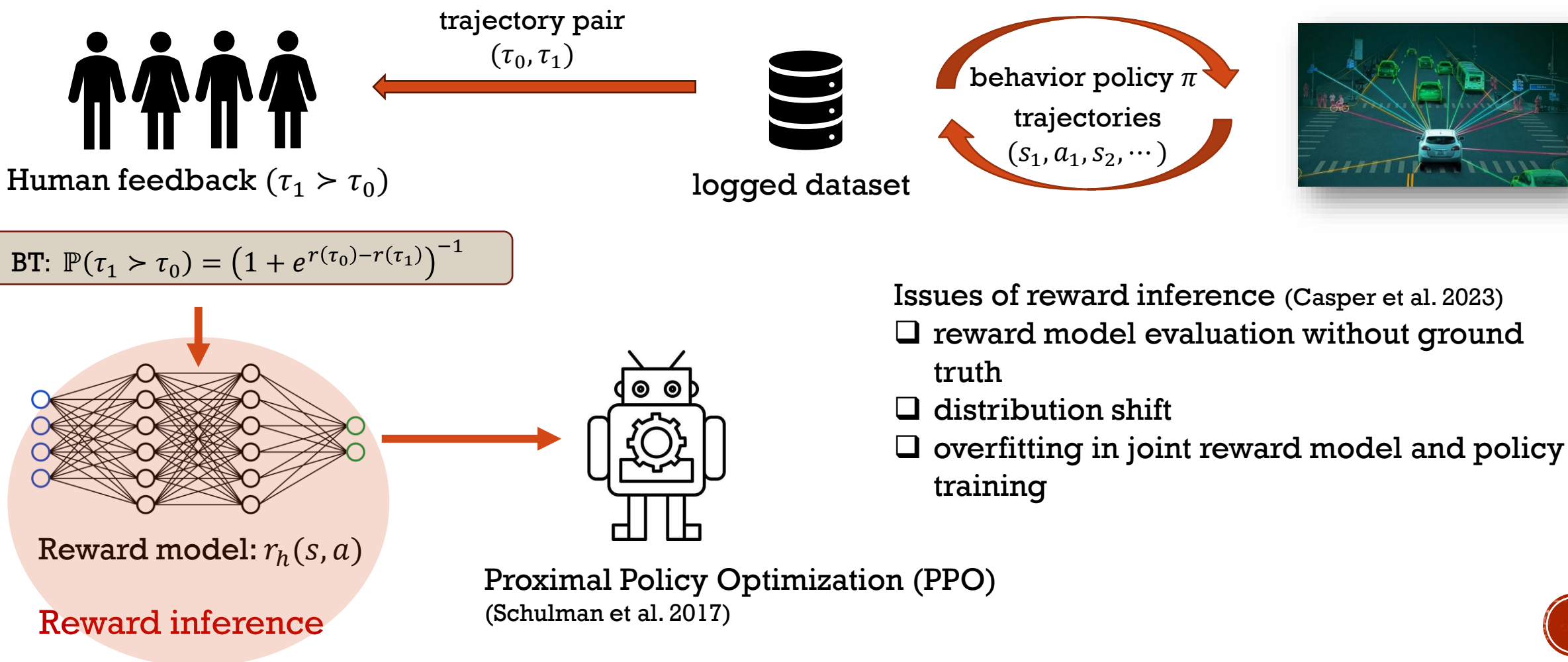
Proximal Policy Optimization (PPO)
(Schulman et al. 2017)

Theoretical analysis of RLHF

- ☐ Value-based (Wang, Liu&Jin 2023, ...)
- ☐ Policy-optimization-based (Du et al. 2024, ...)

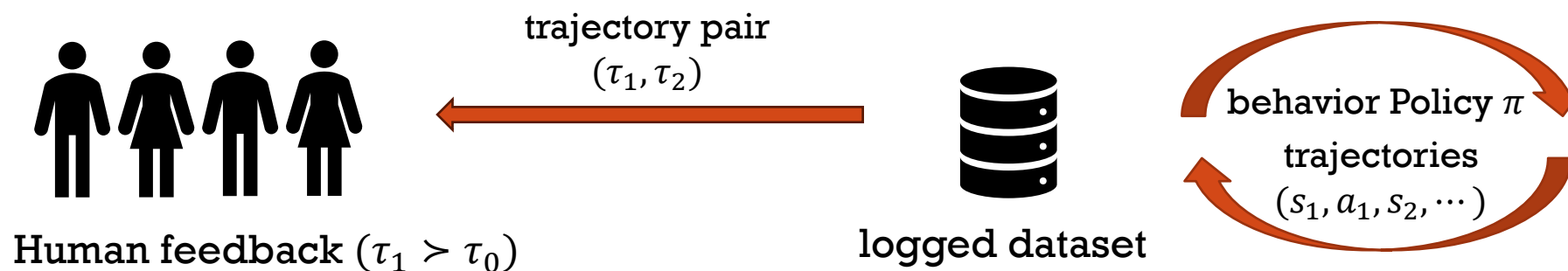
RLHF: The GPT-4 Approach

- Bradley-Terry (BT) Preference Model + Reward Inference + PPO (Achiam et al. 2023)



Direct Preference Optimization

■ DPO (Rafailov et al. 2023)



human preference



reward



optimal policy

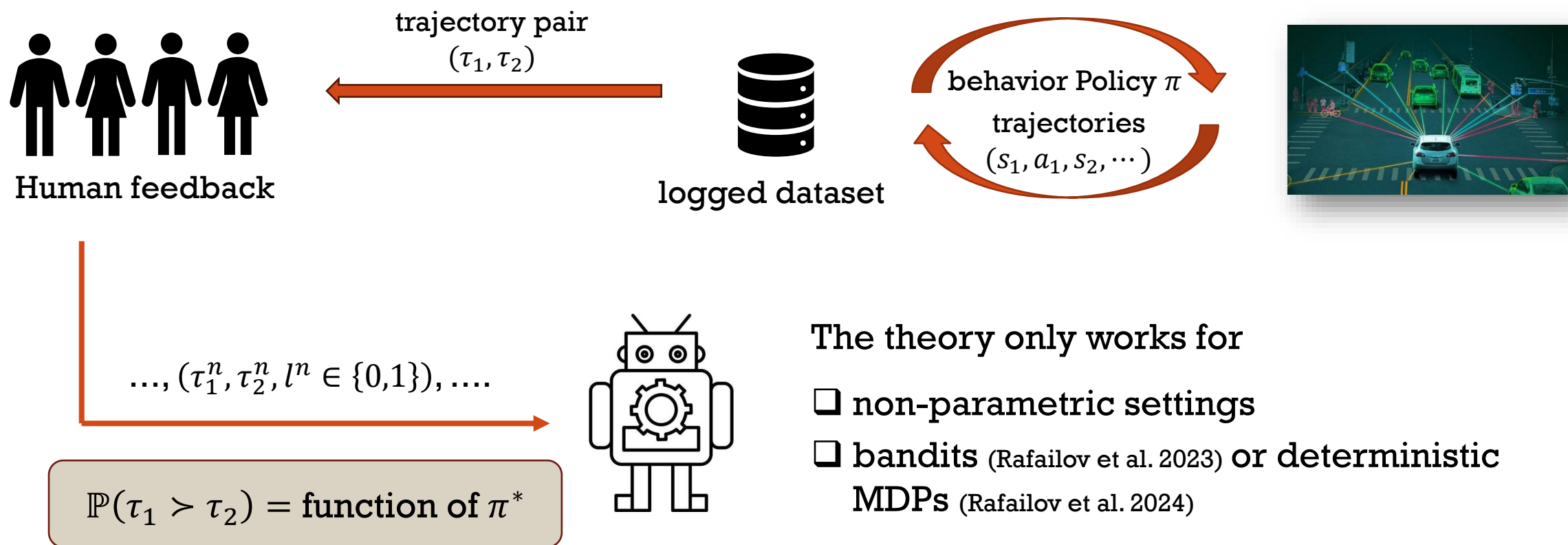
BT model

$$\mathbb{P}(\tau_1 > \tau_0) = (1 + e^{r(\tau_0) - r(\tau_1)})^{-1}$$

nonparametric; bandit setting

$$\pi^*(a|x) \propto \pi_{\text{ref}}(a|x) \exp\left(\frac{1}{\beta} r(x, a)\right)$$

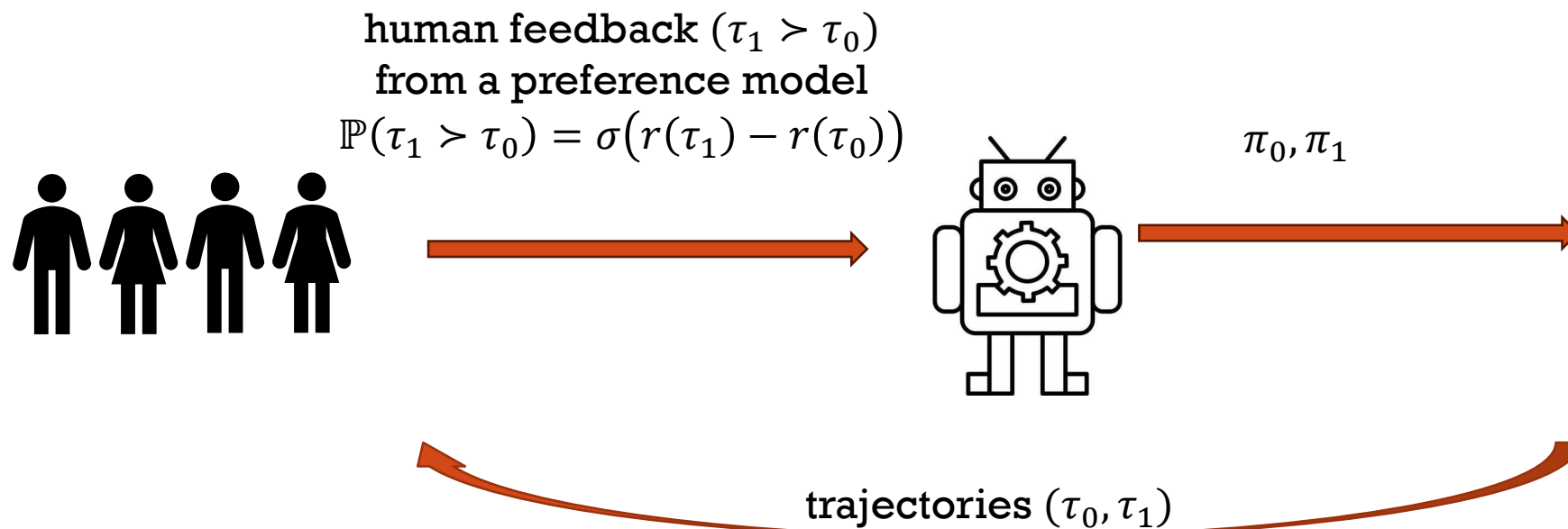
DPO



The theory only works for

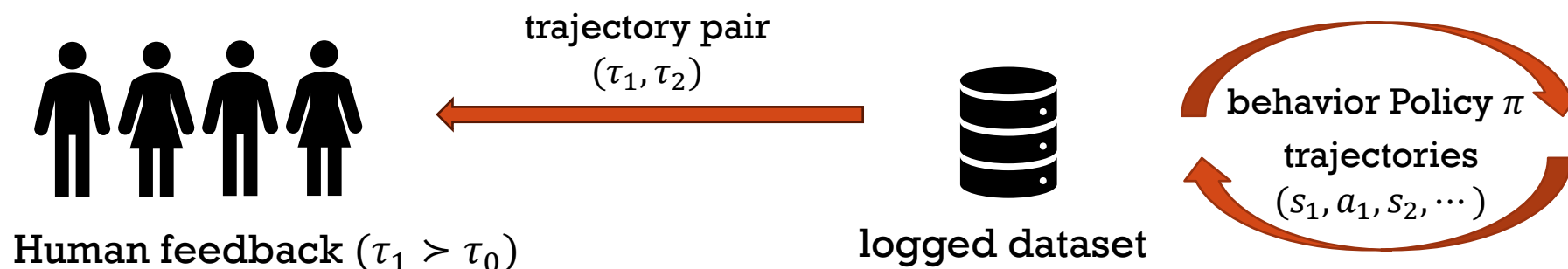
- ❑ non-parametric settings
- ❑ bandits (Rafailov et al. 2023) or deterministic MDPs (Rafailov et al. 2024)

RLHF without Reward Inference for General RL Problems



General RL, Parameterized Policy, Uncountable (s,a)

■ DPO (Rafailov et al. 2023)



human preference



reward



~~optimal policy~~

BT model

$$\mathbb{P}(\tau_1 > \tau_0) = (1 + e^{r(\tau_0) - r(\tau_1)})^{-1}$$

nonparametric; bandit setting

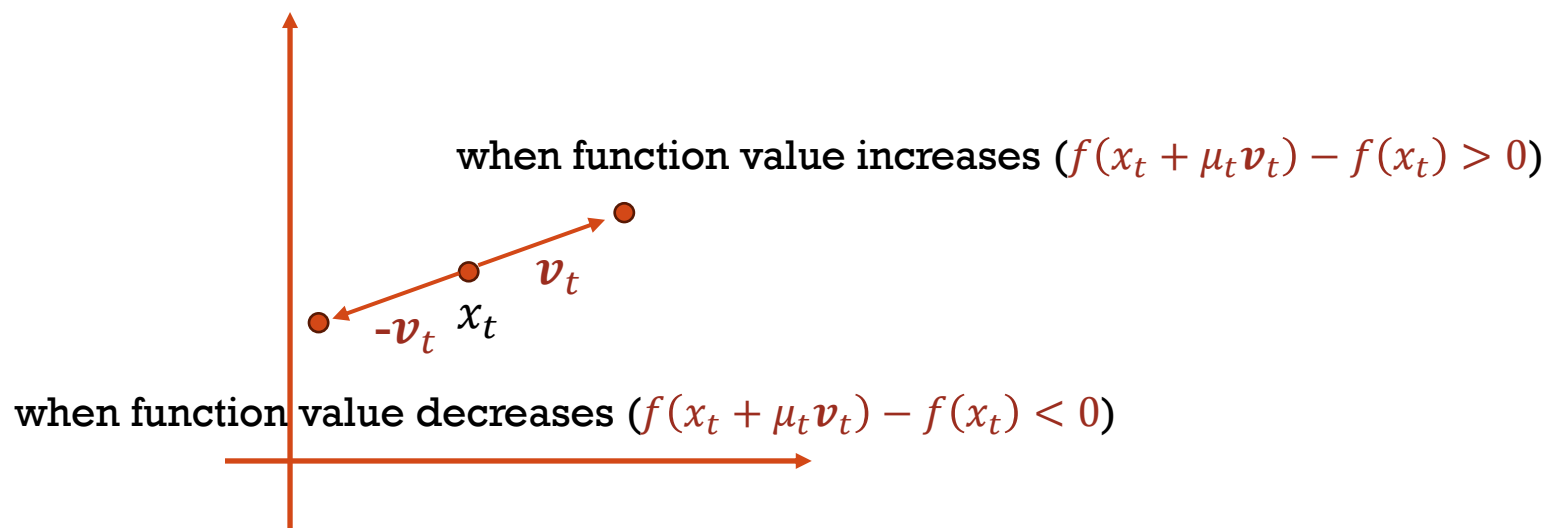
~~$$\pi^*(a|x) \propto \pi_{\text{ref}}(a|x) \exp\left(\frac{1}{\beta} r(x, a)\right)$$~~

Zeroth-Order Optimization

- Two-point, zeroth-order optimization (Nesterov&Spokoiny, 2017)

$$\max f(x) \quad x_{t+1} = x_t + \frac{h_t}{\mu_t} (f(x_t + \mu_t \mathbf{v}_t) - f(x_t)) \mathbf{v}_t$$

- $\mathbf{v}_t$: sampled uniformly from a unit sphere



Zeroth-Order Optimization

- Two-point, zeroth-order optimization
(Nesterov&Spokoiny, 2017)

$$\max f(x)$$

$$x_{t+1} = x_t + \frac{h_t}{\mu_t} (f(x_t + \mu_t \mathbf{v}_t) - f(x_t)) \mathbf{v}_t$$

- $\mathbf{v}_t$: sampled uniformly from a unit sphere

Policy optimization in RL

$$\max_{\theta} V(\pi_{\theta})$$

Value function $V(\pi_{\theta_t})$ under policy π_{θ_t}

Parameter θ_t that defines policy π_{θ_t}

Zeroth-Order Optimization

- Two-point, zeroth-order optimization
(Nesterov&Spokoiny, 2017)

$$\max f(x)$$
$$x_{t+1} = x_t + \frac{h_t}{\mu_t} (f(x_t + \mu_t v_t) - f(x_t)) v_t$$

- v_t : sampled uniformly from a unit sphere

Policy optimization in RL

$$\max_{\theta} V(\pi_{\theta})$$

Value function $V(\pi_{\theta_t})$ under policy π_{θ_t}

Parameter θ_t that defines policy π_{θ_t}

How to obtain the value function difference?

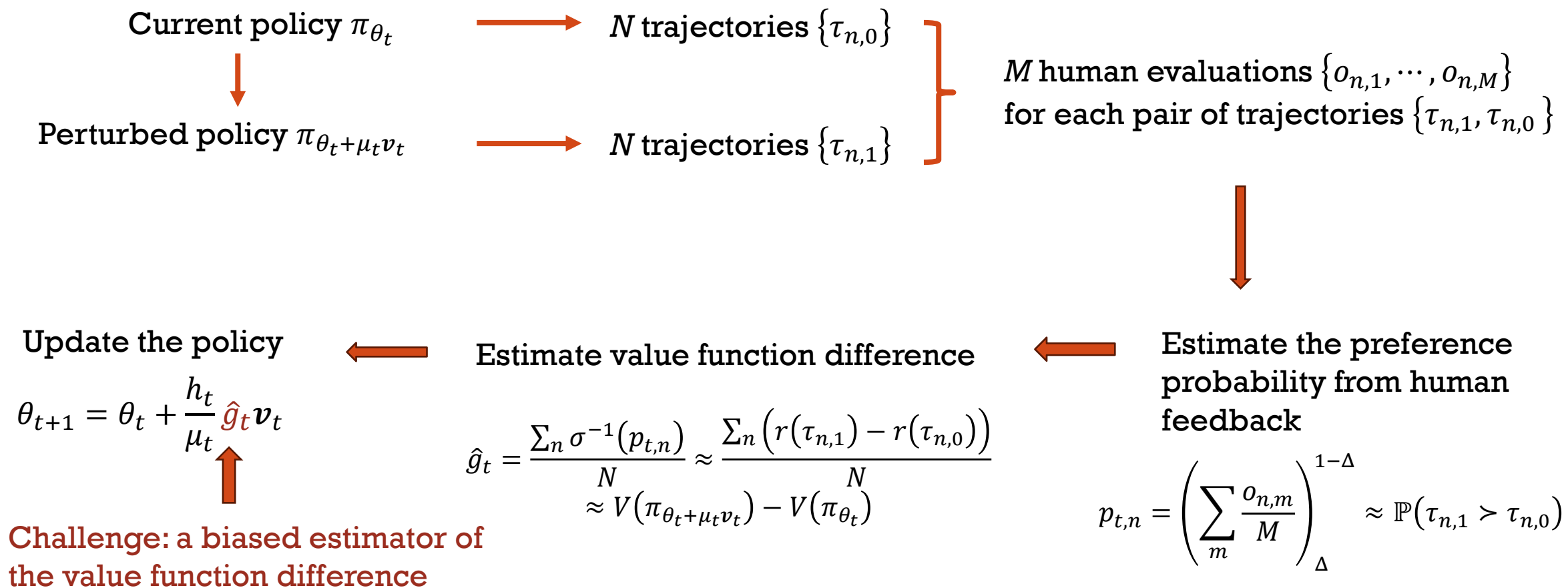
Human feedback!

Zeroth-Order Policy Gradient from Human Feedback

□ ZPG: $\theta_{t+1} = \theta_t + \frac{h_t}{\mu_t} \left(V(\pi_{\theta_t + \mu_t v_t}) - V(\pi_{\theta_t}) \right) v_t$

□ ZPG in HF: $\theta_{t+1} = \theta_t + \frac{h_t}{\mu_t} \hat{g}_t v_t$

Zeroth-Order Policy Gradient from Human Feedback



Rate of Convergence

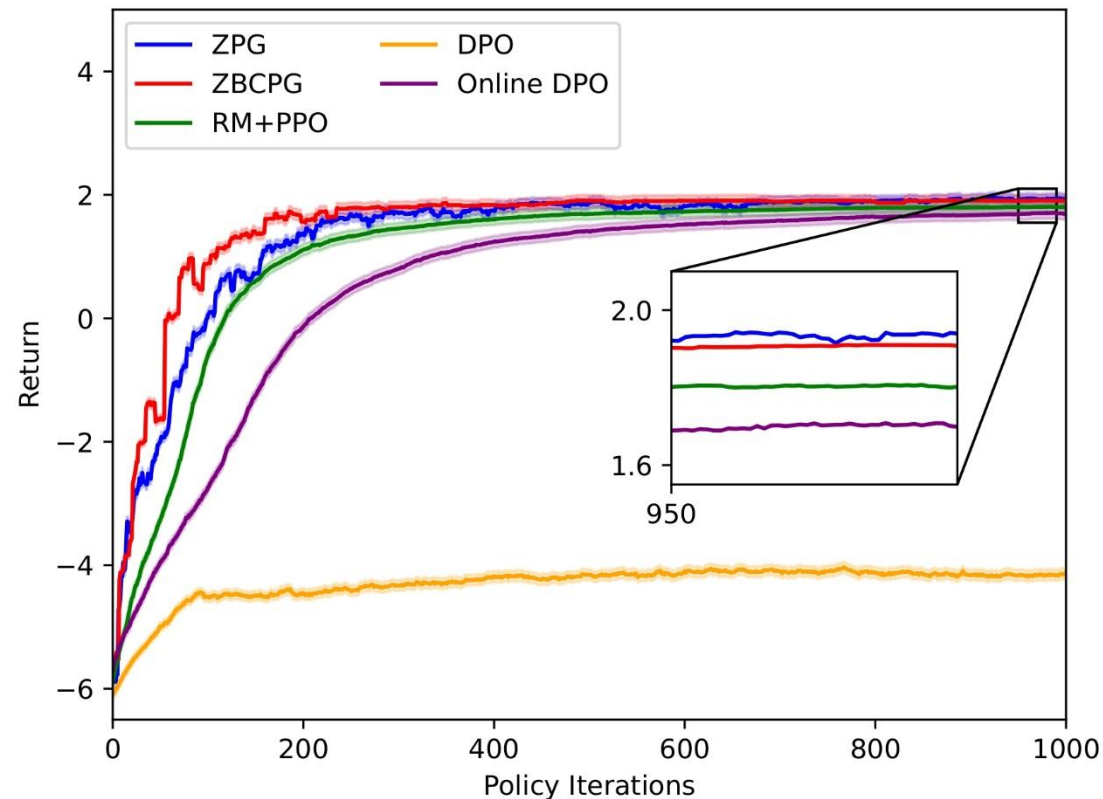
Theorem: ZPG has the following rate of convergence to a stationary point:

$$\tilde{O}\left(\frac{Hd}{T} + \frac{d^2\sqrt{\log M}}{\sqrt{M}} + \frac{Hd\sqrt{d}}{\sqrt{N}}\right)$$

- d : policy dimension
- T : number of policy gradient steps
- N : number of trajectory pairs per pg step
- M : number of human preferences per trajectory pair

Experiments

- Stochastic Gridworld
 - Softmax policy parameterization
 - Bradley-Terry human feedback
 - Weibull human feedback



Thanks!
Questions?