



Online Distributionally Robust Reinforcement Learning with **General Function Approximation**



Dr. Debamita Ghosh

Postdoctoral Researcher

Department of Electrical and Computer Engineering

University of Central Florida, FL

Prof. George K. Atia, *Department of Electrical and Computer Engineering, University of Central Florida, FL*

Prof. Yue Wang, *Department of Electrical and Computer Engineering, University of Central Florida, FL*

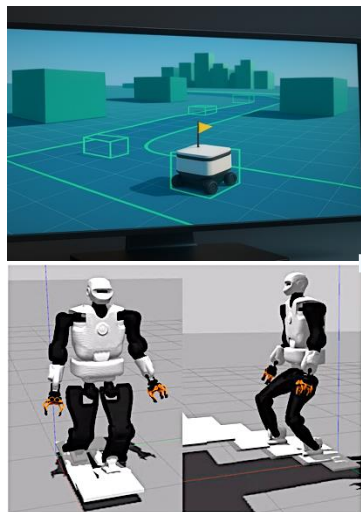
ICML Poster Presentation
Seoul, South Korea
July 6-11, 2026



ICML

Motivation: Bridge Sim-to-Real Gap

Simulation



Real-World Deployment



IHMC's Atlas robot **failed** on the **first day** of the **DARPA Robotics Challenge Finals**.

nominal training world (π trained on P^*)



train and deploy



Sim-to-Real Gap



shifted deployment world
(P^* may shift)

Illustration of a practical problem. A detector trained only on synthetic images makes errors (red crosses) when deployed in real scenes due to the sim-to-real gap.

Sim-to-real gap. Real deployments face unpredictable **noise**, unmodelled **perturbations** and even **adversarial attacks** that are **not** fully captured by the simulator.

*Policies that look excellent in simulation can **fail catastrophically** after deployment.*

Vanilla RL

Optimizes expected return under the training dynamics only.

Distributionally Robust RL (DR-RL)

Optimizes performance under a divergence ball around the nominal dynamics.

Online challenge

Data come from P^* , while robust value depends on worst-case dynamics.

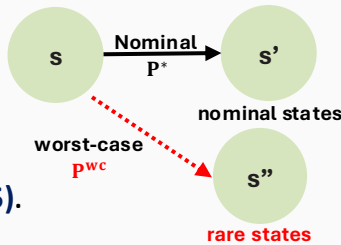
Key Challenges: Why is Online Robust RL difficult?

Purely **online** robust reinforcement learning with **general function approximation** is fundamentally hard.

Off-Dynamics Learning

Data come from **nominal dynamics**.

Robustness is evaluated under **worst-case dynamics**.



Key Challenge:

- **Distribution mismatch**
 - **Information bottleneck**
- (Lu et al. 2024, He et al. 2025).

Existing methods need generative model or rich offline data or strong coverage assumptions (Shi et al. 2023; Shi and Chi 2024).

Robust Bellman Learning

Robust Bellman update requires optimization over uncertainty sets.

$$\left(\mathcal{T}_h^{\phi, \sigma} V\right)(s, a) = r(s, a) + \inf_{P \in \mathcal{U}_{\phi, \sigma}} E_P[V]$$

Key Challenge:

- Worst-case operator is **not directly observable**.
- **Inner optimization** required at every (s, a) .

Existing Standard fitted RL/UCB assume nominal Bellman updates (Jin et al., 2021). Robust RL requires learning a worst-case Bellman operator online (Panaganti et al., 2022; 2024).

Scalability

Real-world robust RL requires:

- **neural networks**
- **continuous states**
- **large action spaces**

Key Challenge:

- **Tabular bonuses** scale with: $|\mathcal{S}| |\mathcal{A}|$.
- **Linear function approximation methods** require restrictive linearly structure assumptions.

No general intrinsic-complexity theory existed for online DR-RL.

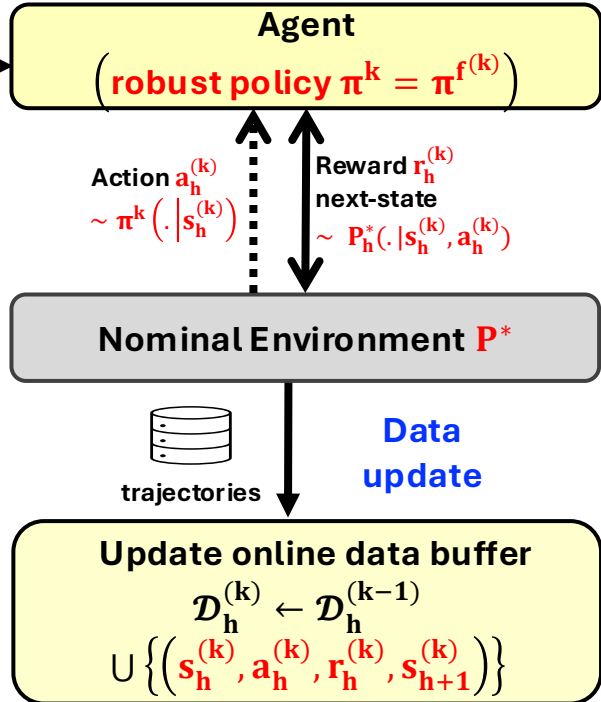
Can we learn a robust policy purely online, with general function approximation, and provable regret?

- ❑ Shi, L.; Li, G.; Wei, Y.; Chen, Y.; Geist, M.; and Chi, Y. 2023. The Curious Price of Distributional Robustness in Reinforcement Learning with a Generative Model. *Advances in Neural Information Processing Systems*, 36: 79903–79917.
- ❑ Shi, L.; and Chi, Y. 2024. Distributionally Robust Model-Based Offline Reinforcement Learning with Near-Optimal Sample Complexity. *Journal of Machine Learning Research*, 25(200): 1–91.
- ❑ Lu, M.; Zhong, H.; Zhang, T.; and Blanchet, J. 2024. Distributionally Robust Reinforcement Learning with Interactive Data Collection: Fundamental Hardness and Near-Optimal Algorithm. *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- ❑ He, Y.; Liu, Z.; Wang, W.; and Xu, P. 2025. Sample Complexity of Distributionally Robust Off-Dynamics Reinforcement Learning with Online Interaction. *In Forty-second International Conference on Machine Learning*.
- ❑ Jin, C., Liu, Q., and Miryoosefi, S. Bellman Eluder Dimension: New Rich Classes of RL Problems, and Sample-Efficient Algorithms. *Advances in neural information processing systems*, 34:13406–13418, 2021.
- ❑ Panaganti, K., Xu, Z., Kalathil, D., and Ghavamzadeh, M. Robust Reinforcement Learning using Offline Data. *Advances in neural information processing systems*, 35:32211–32224, 2022.
- ❑ Panaganti, K., Wierman, A., and Mazumdar, E. Model-Free Robust ϕ -Divergence Reinforcement Learning Using Both Offline and Online Data. *arXiv preprint arXiv:2405.05468*, 2024.

Algorithm: Robust Fitted Learning with ϕ -Divergence (RFL- ϕ)

- $\mathcal{F}$: value functional class
- $\mathcal{G}$: dual functional class
- β : confidence parameter
- σ : robustness radius

Interaction with Nominal Env. & Data update



Dual Robust Bellman Learning (Our Core)

1. Learn dual function (for any $f_{h+1} \in \mathcal{F}_{h+1}^{(k)}$)

$$\underline{g}_{f_{h+1}}^{(k)} = \operatorname{argmin}_{g \in \mathcal{G}} \widehat{\text{DualLoss}}_h^{(k)}(g; f_{h+1})$$

2. Approximate robust Bellman backup

$$\hat{\mathcal{T}}_{\underline{g}_{f_{h+1}}^{(k)}}^{\phi, \sigma} f_{h+1} \text{ (dual-based robust Bellman operator)}$$

3. Compute loss for candidate f_h

$$\mathcal{L}_h^{(k)}(f_h, f_{h+1}, \underline{g}_{f_{h+1}}^{(k)}) \text{ -- robust fitting loss on } \mathcal{D}_h^{(k)}$$

Confidence set over functional class $\mathcal{F}$

Build confidence set for step h:

$$\mathcal{F}_h^{(k)} = \left\{ f_h \in \mathcal{F}_h : \begin{aligned} &\mathcal{L}_h^{(k)}(f_h, f_{h+1}, \underline{g}_{f_{h+1}}^{(k)}) \\ &- \min_{f'_h \in \mathcal{F}_h} \mathcal{L}_h^{(k)}(f'_h, f_{h+1}, \underline{g}_{f_{h+1}}^{(k)}) \leq \beta, \forall h \in [H] \end{aligned} \right\}$$

Optimistic Robust Planning & Greedy Policy

Optimistic robust value estimation

$$f^{(k)} = \operatorname{argmax}_{f \in \mathcal{F}^{(k-1)}} f_1(s_1, \pi_1^f(s_1))$$

Greedy robust policy

$$\pi^{(k)} = \pi^{f^{(k)}}$$

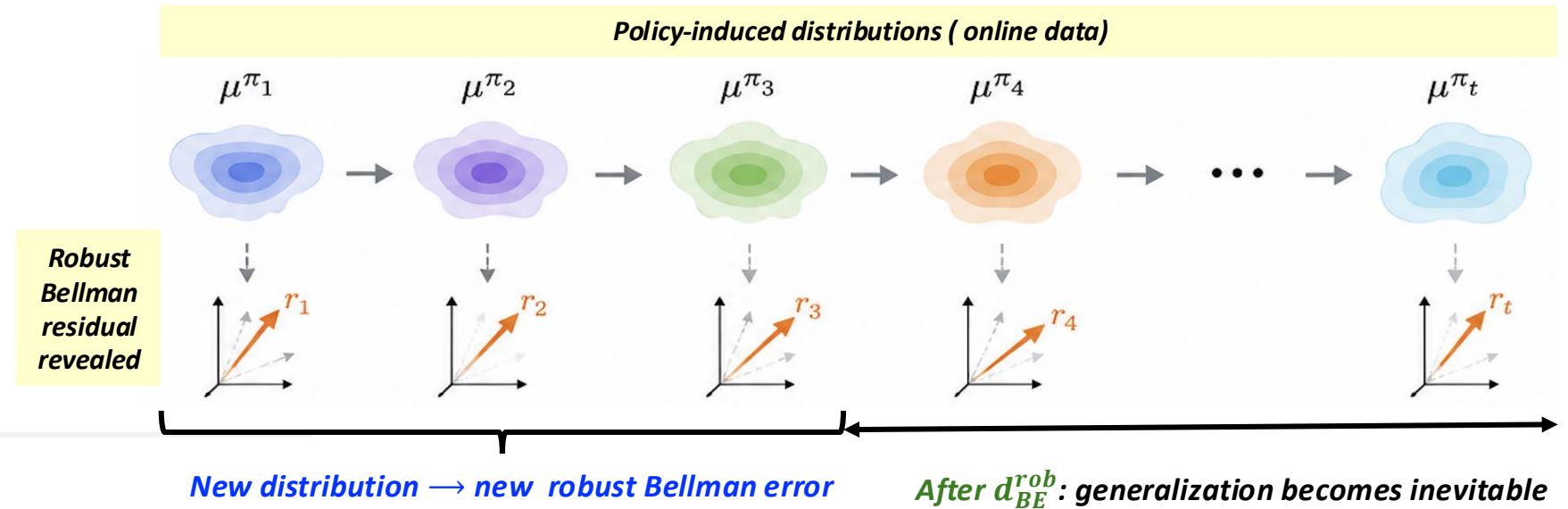
- Choose the **most optimistic value function** in the **confidence set** and **act greedily**.

Repeat over episodes k

Theoretical Guarantee: Why this matters ?

Robust Bellman-Eluder (BE) dimension: $d_{\text{BE}}^{\text{rob}} = \underbrace{\dim_{\text{DE}}}_{\text{Eluder dimension}} \left(\underbrace{\left(\mathcal{I} - \mathcal{T}_h^{\phi, \sigma} \right) \mathcal{F}}_{\text{Robust Bellman residual class}} \right)$

- intrinsic **exploration** complexity under robust Bellman updates.
- Counts how many **policy distributions** reveal **new, statistically, independent** Bellman errors.



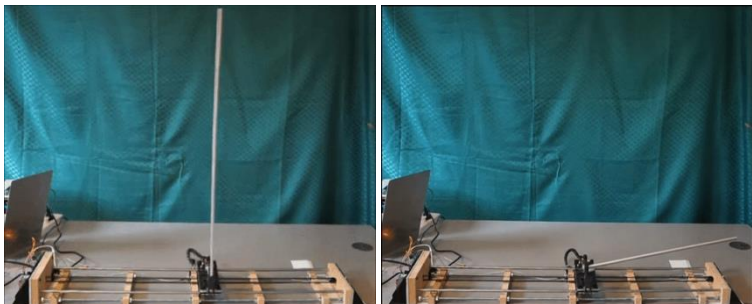
Why it matters?

- Smaller $d_{\text{BE}}^{\text{rob}} \Rightarrow$ fewer essentially different robust Bellman errors.
- Exploration scales with $d_{\text{BE}}^{\text{rob}}$ not $|\mathcal{S}| |\mathcal{A}| \Rightarrow$ **sample efficient**.
- First regret guarantee for purely online DR-RL with general function approximation:

$$\text{Regret}(K) \leq \tilde{O} \left(\sqrt{d_{\text{BE}}^{\text{rob}} H^2 B_{\phi}^2(\sigma) K} + \epsilon^{\text{dual}} \right).$$

- Recovers near-optimal tabular and linear robust RL guarantees.

Experiments: Cartpole



CartPole policies trained nominally, then tested under action and dynamics perturbations.

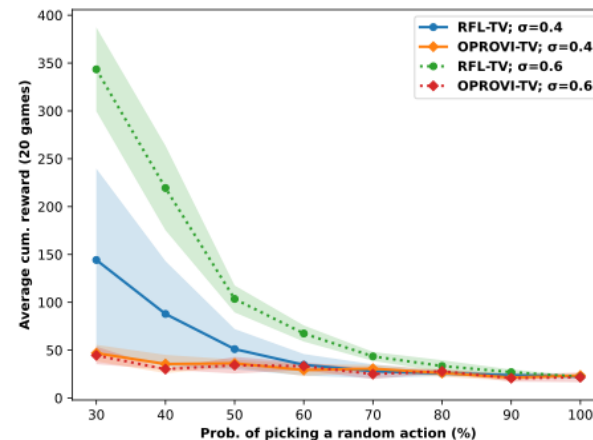
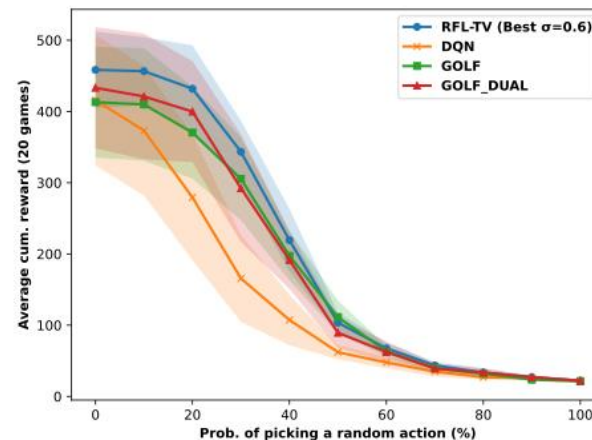
Strong robustness under perturbations

- Up to 2 × higher reward than DQN; competitive with tabular OPROVI-TV.

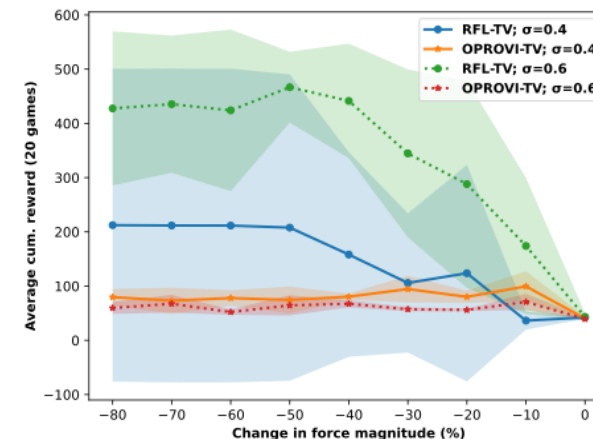
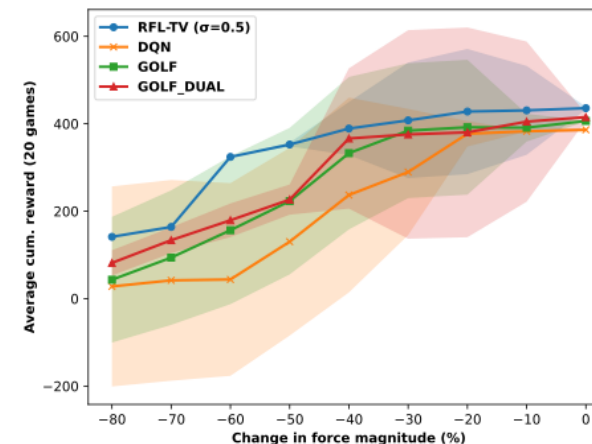
Stable under severe dynamics shifts

- Robust even under $-80\% \sim -40\%$ force perturbations; non-robust FA baselines collapse.

Action
Perturbation



Force-magnitude
Perturbation



RFL-TV vs. Functional
Approximation (FA) Algorithms

RFL-TV vs. OPROVI-TV (Tabular).

Key Takeaway: Robust fitted Bellman learning achieves strong robustness under severe distribution shifts without relying on tabular representations.

DARING TO
BOLDLY INVENT



FUTURE

THANK YOU!

ANY QUESTIONS?



de881780@uc
f.edu



yue.wang@ucf
.edu



george.at
ia@ucf.edu



Poster Q&A & Follow-Up:

<https://forms.gle/FRMRDh4Toj2GY5wC7>