

Robust Reinforcement Learning as a Stackelberg Game via Adaptively-Regularized Adversarial Training

Peide Huang*, Mengdi Xu, Fei Fang and Ding Zhao

Carnegie Mellon University

{peideh, mengdixu}@andrew.cmu.edu, feif@cs.cmu.edu, dingzhao@cmu.edu

Abstract

Robust Reinforcement Learning (RL) focuses on improving performances under model errors or adversarial attacks, which facilitates the real-life deployment of RL agents. Robust Adversarial Reinforcement Learning (RARL) is one of the most popular frameworks for robust RL. However, most of the existing literature models RARL as a zero-sum simultaneous game with Nash equilibrium as the solution concept, which could overlook the sequential nature of RL deployments, produce overly conservative agents, and induce training instability. In this paper, we introduce a novel sequential formulation of robust RL — a general-sum Stackelberg game model called RRL-Stack — to formalize the sequential nature and provide extra flexibility for robust training. We develop the Stackelberg Policy Gradient algorithm to solve RRL-Stack, leveraging the Stackelberg learning dynamics by considering the adversary’s response. Our method generates challenging yet solvable adversarial environments which benefit RL agents’ robust learning. Our algorithm demonstrates better training stability and robustness against different testing conditions in the single-agent robotics control and multi-agent high-way merging tasks.

1 Introduction

Deep reinforcement learning (DRL) has demonstrated great potential in handling complex tasks. However, its real-life deployments are hampered by the commonly existing discrepancies between training and testing environments, e.g., uncertain physical parameters in robotics manipulation tasks [Zhao *et al.*, 2020; Xu *et al.*, 2020] and changing hidden strategies of surrounding vehicles in autonomous driving scenarios [Ding *et al.*, 2021; Xu *et al.*, 2021]. To remedy the fragility against model mismatches, recent advances in robust reinforcement learning (RRL) [Morimoto and Doya, 2005] propose to learn robust policies that maximize the worst-case performances over various uncertainties.

*Contact Author

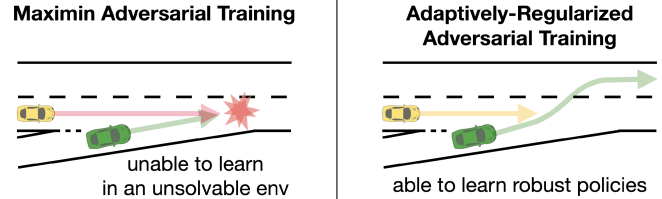


Figure 1: A high-level comparison between the existing RARL formulation and our RRL-Stack formulation for robust RL. In RARL, the RL agent (green car) is trained with an adversary (yellow car) that generates extremely challenging and even unsolvable environments. In RRL-Stack, the RL agent is trained with an adaptively-regularized adversary that generates challenging yet solvable environments to improve robustness against different testing environments.

One popular RRL framework is the Robust Adversarial Reinforcement Learning (RARL) [Pinto *et al.*, 2017], which treats environment mismatches as adversarial perturbations against the agent. RARL formulates a two-player zero-sum simultaneous game between the *protagonist* who aims to find a robust strategy across environments and the *adversary* who exerts perturbations. Computational methods have been proposed to solve this game and find a robust strategy for the protagonist. Despite promising empirical performances in many tasks, existing approaches under the RARL framework have three potential limitations: 1) overlook the sequential nature in the deployments of RL agents, 2) produce overly conservative agents, and 3) induce training instability. These limitations will be discussed in detail in Section 3.2.

In this paper, we propose a novel robust RL formulation — *Robust Reinforcement Learning as a Stackelberg Game (RRL-Stack)* to address these limitations. To formalize the sequential structure that the protagonist is trained first and then deployed in uncertain environments, we model robust RL as a Stackelberg game where the protagonist is the leader, and the adversary is the follower. By assuming the adversary optimizes a linearly combined objective of two extreme scenarios for the protagonist, RRL-Stack enables the protagonist to learn a robust policy in challenging yet solvable adversarial environments. RRL-Stack further provides extra flexibility to control the protagonist’s conservativeness or accommodate more general multi-objective training settings. We then leverage the Stackelberg learning [Fiez *et al.*, 2020] and develop the Stackelberg Policy Gradient (Stack-PG) method which has known

convergence and stability guarantees to find the local surrogate solutions of the RRL-Stack.

In summary, our main contributions are: a) introducing RRL-Stack, a novel Stackelberg game-theoretical formulation of RRL (Section 4.1), b) developing a Stack-PG, a policy gradient method for solving the game (Section 4.2), c) demonstrating that RRL-Stack formulation together with Stack-PG reduces the training instability significantly compared with existing methods, and agents learn robust but not overly conservative policies from challenging yet solvable adversarial environments (Section 5).

2 Preliminaries

Markov Decision Process. We consider a Markov Decision Process (MDP) defined by the 5-tuple $(\mathcal{S}, \mathcal{A}, \mathcal{P}, r, \gamma, \mu_0)$, where $\mathcal{S}$ is a set of states, $\mathcal{A}$ is a set of continuous or discrete actions. $\mathcal{P} : \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ is the transition probability ($\Delta(\mathcal{S})$ is the distribution over $\mathcal{S}$), $r : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ is the reward function, γ is the discount factor and μ_0 is the initial state distribution. The goal of RL is to find a policy $\pi_\theta : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ parametrized by θ that maximizes the expected return $\mathbb{E}_{\tau \sim \pi_\theta} [R(\tau)] := \mathbb{E}_{\tau \sim \pi_\theta} \left[\sum_{t=0}^{T-1} \gamma^t r(s_t, a_t) \right]$, where τ denotes the trajectories sampled using the policy π_θ . For simplicity, we overload the symbol for policies as the policy parameters $\mathbb{E}_{\tau \sim \theta} [R(\tau)] := \mathbb{E}_{\tau \sim \pi_\theta} [R(\tau)]$.

Stochastic Game. A two-player stochastic game [Shapley, 1953] is defined by a tuple $(\mathcal{S}, \mathcal{A}_1, \mathcal{A}_2, \mathcal{P}, r_1, r_2, \gamma, \mu_0)$, where $\mathcal{S}$ is a set of states, $\mathcal{A}_1$ and $\mathcal{A}_2$ are action spaces of Agent 1 and 2, respectively. $\mathcal{P} : \mathcal{S} \times \mathcal{A}_1 \times \mathcal{A}_2 \rightarrow \Delta(\mathcal{S})$ is the transition probability. $r_1, r_2 : \mathcal{S} \times \mathcal{A}_1 \times \mathcal{A}_2 \rightarrow \mathbb{R}$ is the reward function for Agent 1 and Agent 2 respectively. μ_0 is the initial state distribution. If $r_1 = -r_2$, the stochastic game is zero-sum, otherwise general-sum.

3 Related Work

3.1 Robust Reinforcement Learning via Adversarial Training

The transition model in training and testing could sometimes be different in RL applications. For example, in Sim2Real transfer, the physical simulator has unavoidable modeling errors, and the real world always has unexpected uncertainty [Akkaya *et al.*, 2019]. To address this issue, [Pinto *et al.*, 2017] proposes Robust Adversarial Reinforcement Learning (RARL), which introduces an adversary to apply perturbations to the environment dynamics, such as disturbing forces applied to the robot’s joints, gravity constants, and friction coefficients.

RARL formulates a simultaneous zero-sum stochastic game. More concretely, let θ be the policy parameters of the protagonist who acts in an uncertain environment and ψ be the policy parameters of the adversary who controls the environment dynamics. RARL assumes that the protagonist maximizes $\mathbb{E}_{\tau \sim \theta, \psi} [R(\tau)]$ while the adversary minimizes it. [Tessler *et al.*, 2019] introduces Noisy Robust MDP to apply perturbations to the commanded action without access to the environment simulator. [Kamalaruban *et al.*, 2020] adapts Langevin Learning Dynamics to approach robust RL from a sampling

perspective. In addition, most existing works use gradient-descent-ascent-based algorithms to train agents.

3.2 Limitations of Existing Methods

Despite the empirical success in many tasks, there are a few limitations in the existing formulation.

Overlooking the sequential nature. Most of the existing methods use Nash equilibrium (NE) as the solution concept, which means no agent can achieve higher rewards by deviating from the equilibrium strategy. While NE is a standard solution concept for simultaneous-move games, using NE for RRL overlooks the sequential order of actual policy deployments, as the protagonist’s policy is chosen first and then tested in a range of uncertain environments [Jin *et al.*, 2020]. Considering this sequential nature, a more appropriate solution concept is Stackelberg Equilibrium (SE). Although in the many zero-sum games, NE also coincides with SE, it is not the case in deep RARL as discussed by [Tessler *et al.*, 2019; Zhang *et al.*, 2019]. One of the reasons is that with policy parametrizations in Deep RL, e.g., using neural networks, the objectives of stochastic games are nonconvex-nonconcave in the parameter space, and as a result, maximin value is not equal to minimax value in general. Different from most existing works, our RRL-Stack adopts a sequential game-theoretical formulation and solution concept, which formalizes the sequential order of actual policy deployments.

Producing overly conservative agents. Existing work often assumes that the protagonist and adversary play a zero-sum game where the adversary minimizes the protagonist’s expected return. However, such a formulation encourages the adversary to generate extremely difficult, even unsolvable environments for the protagonist. As a result, the protagonist may choose an overly conservative strategy or even not be able to learn any meaningful policies because the most adversarial environment is completely unsolvable. [Dennis *et al.*, 2020] proposes PAIRED to mitigate this issue by replacing the return in the objective with the minimax regret, which has the closest relation to our proposed method. The protagonist’s regret is defined as the difference between the maximum possible return and the current return, which can be seen as a special case of our adversary’s objective. Some other literature [Shen *et al.*, 2020] dealing with state-robustness uses l_p ball to constrain the adversarial perturbation on the state.

Potential training instability. Besides the two aforementioned issues related to the game formulation, the gradient-descent-ascent learning dynamics in RARL could lead to significant training instability even in the simple linear-quadratic system. The unstable training is partially due to the non-stationarity of environments, as the environment controlled by the adversary could be constantly evolving. Recent research in RARL [Zhang *et al.*, 2020; Yu *et al.*, 2021] aims to tackle this problem, yet only for linear-quadratic systems with the zero-sum formulation. In contrast, our algorithm relying on the Stackelberg learning dynamics [Fiez *et al.*, 2020] can be applied to more general RL settings.

4 Method

We first introduce our proposed formulation, Robust Reinforcement Learning as a Stackelberg Game (RRL-Stack) in Section 4.1. We then present our main algorithm, Stackelberg Policy Gradient (Stack-PG) in Section 4.2.

4.1 Robust RL as a Stackelberg Game

Although formulated as simultaneous games in most of the existing works, adversarial and robust training are in fact sequential [Jin *et al.*, 2020]. In robust RL via adversarial training, the protagonist has to choose its policy first and then the adversary chooses the best responding policy given the protagonist’s policy. To formalize this inherently sequential structure, we formulate robust RL as a Stackelberg game as follows:

$$\begin{aligned} \max_{\theta \in \Theta} \mathbb{E}_{\tau \sim \theta, \psi} [R_{\text{pro}}(\tau)] \quad & \text{s.t.} \\ \psi \in \arg \max_{\psi' \in \Psi} \mathbb{E}_{\tau \sim \theta, \psi'} [-R_{\text{pro}}(\tau)], \end{aligned}$$

where θ and ψ parametrize the protagonist’s and the adversary’s policy, respectively. $\mathbb{E}_{\tau \sim \theta, \psi} [R_{\text{pro}}(\tau)]$ is the expected return of the protagonist.

With the sequential game structure, we now aim to address the limitations caused by the widely adapted zero-sum formulation. It could result in an overly conservative protagonist’s policy or even disable the protagonist’s effective learning since the adversary is encouraged to produce a difficult, even completely unsolvable environment. To address this limitation, we introduce an oracle term that encodes the highest possible return of the protagonist in the current adversarial environment. Inspired by the alpha-maxmin expected utility in the economics literature [Li *et al.*, 2019], we linearly combined this oracle term with the original adversary’s objective. Formally,

$$\max_{\theta \in \Theta} \mathbb{E}_{\tau \sim \theta, \psi} [R_{\text{pro}}(\tau)] \quad \text{s.t.} \quad (1)$$

$$\psi \in \arg \max_{\psi' \in \Psi} \alpha \mathbb{E}_{\tau \sim \theta, \psi'} [-R_{\text{pro}}(\tau)] + (1 - \alpha)V^*(\psi'), \quad (2)$$

where $V^*(\psi')$ is the highest possible return of the protagonist given the current adversarial environment ψ' . This term is approximated by training an oracle RL agent that is locally optimized. Let the oracle agent’s policy be parametrized by ω , $V^*(\psi') := \sup_{\omega} \mathbb{E}_{\tau \sim \omega, \psi'} [R_{\text{ora}}(\tau)]$. This term can be seen as a regularization of the adversary’s objective. The adversary is incentivized to generate environments that are challenging to solve for the current protagonist but solvable for the oracle agent. As the protagonist learns to solve the current environment, the adversary is forced to find harder environments to receive a higher reward, adaptively increasing the difficulty of the generated environments.

The coefficient $\alpha \in [0, 1]$ balances how adversarial the environment is and linearly combines two extreme scenarios:

- When $\alpha = 1$, RRL-Stack’s solution corresponds to the *Maximin* robust strategy.
- When $\alpha = 0.5$, the solution corresponds to the strategy when the adversary maximizes the protagonist’s regret, which is defined by the difference between the maximum possible return and the current return.

- When $\alpha = 0$, the solution corresponds to the *Maximax* strategy, where the protagonist chooses a strategy that yields the best outcome in the most optimistic setting.

Denote $f_{\text{pro}}(\theta, \psi) := \mathbb{E}_{\tau \sim \theta, \psi} [R_{\text{pro}}(\tau)]$ and $f_{\text{adv}}(\theta, \psi) := \alpha \mathbb{E}_{\tau \sim \theta, \psi'} [-R_{\text{pro}}(\tau)] + (1 - \alpha)V^*(\psi')$. The solution to the Stackelberg game is the Stackelberg Equilibrium. We have the following definition:

Definition 4.1 (Stackelberg Equilibrium). *(SE) The joint strategy $(\theta^*, \psi^*) \in \Theta \times \Psi$ of the protagonist is a Stackelberg equilibrium if*

$$\inf_{\psi \in \mathcal{R}(\theta^*)} f_{\text{pro}}(\theta^*, \psi) \geq \inf_{\psi \in \mathcal{R}(\theta)} f_{\text{pro}}(\theta, \psi), \quad \forall \theta \in \Theta,$$

where $\mathcal{R}(\theta) = \{\psi' \in \Psi \mid f_{\text{adv}}(\theta, \psi') \geq f_{\text{adv}}(\theta, \psi), \forall \psi \in \Psi\}$ is the best response set of the adversary, and $\psi^* \in \mathcal{R}(\theta^*)$

We adapt Proposition 4.4 of [Başar and Olsder, 1998] to formalize the relationship between the return at Stackelberg Equilibrium and Nash equilibrium:

Proposition 4.1. *Consider an arbitrary sufficiently smooth two-player general-sum game $(f_{\text{pro}}, f_{\text{adv}})$ on continuous strategy spaces. Let $f_{\text{pro}}^{\text{Nash}}$ denote the supremum of all Nash equilibrium return for the protagonist and $f_{\text{pro}}^{\text{SE}}$ denote an arbitrary Stackelberg equilibrium return for the protagonist. Then, if $\mathcal{R}(\theta)$ is a singleton for every $\theta \in \Theta$, $f_{\text{pro}}^{\text{SE}} \geq f_{\text{pro}}^{\text{Nash}}$.*

However, since one cannot expect to find the global solution of the Stackelberg game efficiently with a general non-convex-non-concave objective, we define the following local equilibrium concept using the sufficient conditions of SE.

Definition 4.2 (Differential Stackelberg Equilibrium). *(DSE) [Fiez et al., 2020] The joint strategy $(\theta^*, \psi^*) \in \Theta \times \Psi$ with $\psi^* = r(\theta^*)$, where r is an implicit mapping defined by $\nabla_{\psi} f_{\text{adv}}(\theta^*, \psi^*) = 0$, is a differential Stackelberg equilibrium if $Df_{\text{pro}}(\theta^*, r(\theta^*)) = 0$ and $\nabla^2 f_{\text{pro}}(\theta^*, r(\theta^*))$ is negative definite ($D(\cdot)$ denotes the total derivative).*

To this end, the key question is how to solve RRL-Stack. We now explain how to develop Stackelberg Policy Gradient (Stack-PG) leveraging the Stackelberg learning dynamics, which have known convergence and stability guarantees to find the DSE under sufficient regularity conditions.

4.2 Stackelberg Policy Gradient (Stack-PG)

Stackelberg learning dynamics. [Fiez *et al.*, 2020] assume that there exists an implicit mapping from θ to the best-response ψ . The protagonist updates its parameters based on the total derivative ($Df_{\text{pro}} := df_{\text{pro}}(\theta, r^*(\theta))/d\theta$) instead of the partial derivative ($\nabla_{\theta} f_{\text{pro}}$). Since the follower chooses the best response $\psi = r^*(\theta)$, the follower’s policy is an implicit function of the leader’s. The leader therefore can take the total derivative of its objective to update its policy:

$$\frac{df_{\text{pro}}(\theta, r^*(\theta))}{d\theta} = \frac{\partial f_{\text{pro}}(\theta, \psi)}{\partial \theta} + \frac{dr^*(\theta)}{d\theta} \frac{\partial f_{\text{pro}}(\theta, \psi)}{\partial \psi} \quad (3)$$

The implicit differentiation term can be computed using the implicit function theorem [Abraham *et al.*, 2012]:

$$\frac{dr^*(\theta)}{d\theta} = \left(\frac{\partial^2 f_{\text{adv}}(\theta, \psi)}{\partial \theta \partial \psi} \right) \left(-\frac{\partial^2 f_{\text{adv}}(\theta, \psi)}{\partial \psi^2} \right)^{-1} \quad (4)$$

Algorithm 1: StackelbergPolicyGradient (Stack-PG)

```

1 Input:  $\{\tau_{\text{pro}}\}^M, \{\tau_{\text{ora}}\}^M, \theta_{k-1}$ , learning rate  $\gamma_\theta$ 
2  $\omega_{S,\theta} \leftarrow \hat{\nabla}_\theta f_{\text{pro}}(\theta, \psi) +$ 
    $\hat{\nabla}_\theta \hat{\nabla}_\psi f_{\text{adv}}(\theta, \psi) \left( -\hat{\nabla}_\psi^2 f_{\text{adv}}(\theta, \psi) + \lambda I \right)^{-1} \hat{\nabla}_\psi f_{\text{pro}}(\theta, \psi)$ 
3  $\theta_k \leftarrow \theta_{k-1} + \gamma_\theta \omega_{S,\theta}$ 
4 Output:  $\theta_k$ 
    
```

Algorithm 2: MultiPolicyGradient

```

1 Input:  $\{\tau_{\text{pro}}\}^M, \{\tau_{\text{ora}}\}^M, \psi_{k-1}$ , learning rate  $\gamma_\psi$ ,
   auto-tuning  $\in \{True, False\}$ , smoothing factor  $\rho$ 
2  $g_1 \leftarrow \frac{\partial}{\partial \psi} \mathbb{E}[-R_{\text{pro}}(\tau)], g_2 \leftarrow \frac{\partial}{\partial \psi} \mathbb{E}_{\tau \sim \omega, \psi}[R_{\text{ora}}(\tau)]$ 
3 if auto-tuning then
4   Find the optimal  $\alpha^*$  by solving (5)
5    $\alpha \leftarrow \rho\alpha + (1 - \rho)\alpha^*$  // Moving average
6 end
7  $\omega_{S,\psi} \leftarrow \alpha g_1 + (1 - \alpha)g_2$ 
8  $\psi_k \leftarrow \psi_{k-1} + \gamma_\psi \omega_{S,\psi}$ 
9 Output:  $\psi_k$ 
    
```

By combining Eq.3 and Eq.4, we obtain the updating rule for the protagonist. The Stackelberg learning dynamics provide local convergence guarantees to DSE under regularity conditions. We include a numerical example in the Appendix B¹ to demonstrate: 1) why SE is a more appropriate solution concept for robust learning than NE, and 2) how the Stackelberg learning dynamics converge to DSE while the gradient-descent-ascent algorithm fails to converge.

The computation of Hessian required by the Stackelberg learning dynamics takes complexity $O(n^2)$, where n is the number of policy parameters. It could be prohibitively slow when n is large. There are some techniques for efficiently computing unbiased Hessian approximations of deep neural networks such as Curvature Propagation [Martens *et al.*, 2012]. The computation of the inverse of Hessian is another burden but can be alleviated by approximation methods such as conjugate gradient [Shewchuk and others, 1994] and minimal residual [Saad and Schultz, 1986]. We leave integrating efficient computation methods to future work.

Protagonist’s updating rule. Based on the Stackelberg learning dynamics, we develop the update rule for the protagonist, Stackelberg Policy Gradient (Stack-PG), as shown in Algorithm 1. Similar to the policy gradient algorithm, we obtain unbiased estimators for the first-order and second-order gradient information based on trajectory samples. Details for the unbiased estimators are in Appendix A. We can also incorporate state-dependent baselines into the gradient estimators to reduce variance. The regularization term λI ensures the Hessian estimate is invertible, where λ is a scalar and I is the identity matrix. Note that as we increase the value of λ , the protagonist’s update first resembles LOLA [Foerster *et al.*, 2017] and eventually becomes the standard policy gradient.

¹<https://arxiv.org/abs/2202.09514>

Algorithm 3: Solving RRL-Stack with Stack-PG

```

1 Input: Protagonist’s policy  $\pi_\theta$ , Adversary’s policy  $\pi_\psi$ ,
   Oracle Agent’s policy  $\pi_\omega$ , Number of trajectories  $M$ 
2 Initialize learnable parameters  $\theta_0, \psi_0, \omega_0$ 
3 for  $k = 1, 2, \dots, N_{\text{iter}}$  do
4    $\{\tau_{\text{pro}}\}^M \leftarrow \text{rollout}(\pi_{\theta_k}, \pi_{\psi_{k-1}})$ 
5    $\{\tau_{\text{ora}}\}^M \leftarrow \text{rollout}(\pi_\omega, \pi_{\psi_{k-1}})$ 
6    $\theta_k \leftarrow$ 
     StackelbergPolicyGradient( $\{\tau_{\text{pro}}\}^M, \{\tau_{\text{ora}}\}^M$ )
7    $\psi_k \leftarrow \text{MultiPolicyGradient}(\{\tau_{\text{pro}}\}^M, \{\tau_{\text{ora}}\}^M)$ 
8   Train  $\pi_\omega$  in the environment given by  $\psi_k$ 
9 end
10 Output:  $\theta_{N_{\text{iter}}}$ 
    
```

Adversary’s updating rule and auto-tuning α . The adversary updates its parameters with a policy-gradient-based algorithm. Since there are two terms in the adversary’s objective, it can be viewed from a multi-objective RL perspective. Instead of manually tuning the value of α , we can dynamically update it using the multiple-gradient descent algorithm (MGDA) [Désidéri, 2012] for multi-objective learning. Let $g_1 := \frac{\partial}{\partial \psi} \mathbb{E}_{\tau \sim \theta, \psi}[-R_{\text{pro}}(\tau)], g_2 := \frac{\partial}{\partial \psi} \mathbb{E}_{\tau \sim \omega, \psi}[R_{\text{ora}}(\tau)]$. We want to find the α^* that approximately maximizes the minimal improvement of the two terms by solving the optimization problem:

$$\min_{\alpha} \frac{1}{2} \|\alpha g_1 + (1 - \alpha)g_2\|^2, \text{ s.t. } \alpha \in [0, 1]. \quad (5)$$

It is the dual form of the primal optimization problem:

$$\max_{\|d\| \leq 1} \min_i \langle d, g_i \rangle, \quad (6)$$

where the optimal $d^* = (\alpha g_1 + (1 - \alpha)g_2)/\lambda$ and λ is the Lagrangian multiplier of the constraints $\|d\| \leq 1$. After solving for α^* , we use exponential moving average to update α smoothly. The updating algorithm of the adversary is shown in Algorithm 2. Note that this α auto-tuning is not required but is a tool for automatic hyper-parameter tuning.

Oracle agent’s updating rule. The oracle agent can be trained using any on- or off-policy optimization algorithm. In practice, we find that performing multiple policy optimization steps for the oracle agent in each iteration usually serves the purpose well.

The Algorithm 3 summarizes our main algorithm. At each iteration, we first rollout trajectories using the protagonist’s and adversary’s policy (Line 4), as well as trajectories using the oracle agent’s and adversary’s policy (Line 5). Next, we use the Stack-PG to update the protagonist’s policy parameters (Line 6). Then we use the policy-gradient-based method to update the adversary’s policy parameters (Line 7). Finally, we train the oracle agent in the current adversarial environment till reaching local convergence (Line 8).

5 Experiments

We conduct experiments to answer the following questions: (Q1) Does our method produce challenging yet solvable environments? (Q2) Does our method improve the robustness and

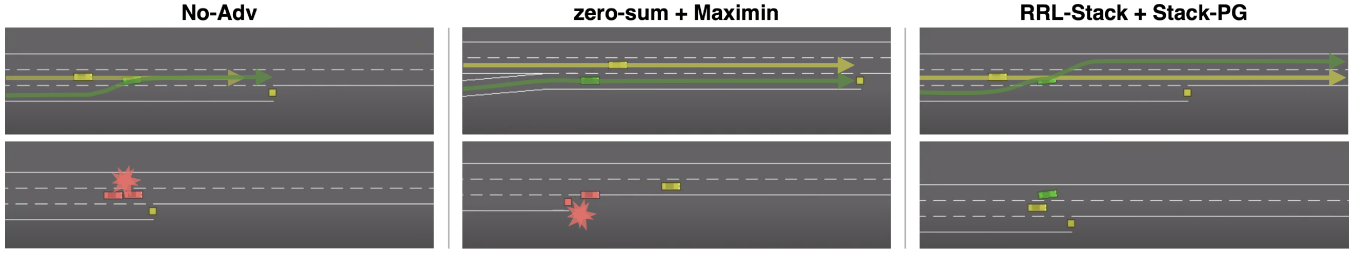


Figure 2: Highway merging policy visualization during testing. The green car is controlled by the protagonist and the yellow car is controlled by the adversary. Cars turning red mean collisions. The opaque lines represent the driving trajectories. The upper row contains moments close to the beginning of the episode, and the lower row contains moments close to the end of the episode.

training stability? **(Q3)** How does the choice of α influence the performance of the protagonist?

5.1 Benchmark Algorithms

We consider 2 game formulations: zero-sum and RRL-Stack, and 3 existing learning algorithms as follows:

Gradient-descent-ascent (GDA). GDA alternates between the policy gradient updates of protagonist and adversary at 1:1 ratio [Zhang *et al.*, 2021].

Maximin operator. Maximin operator is similar to GDA, but the difference is that the adversary updates multiple iterations between each update of the protagonist [Tessler *et al.*, 2019]. In the experiments, we use a ratio of 1:3 to alternate between the updates of the protagonist and the adversary.

Learning with opponent-learning awareness (LOLA). LOLA [Foerster *et al.*, 2017] is a seminal work in considering the opponent while doing gradient ascent. We choose LOLA for its similarity to the Stackelberg learning dynamics.

We consider several combinations of game formulations and learning algorithms: zero-sum game formulation with GDA, Maximin operator, and LOLA; RRL-Stack game formulation with GDA, Maximin operator, and Stack-PG. *Zero-sum + GDA* and *zero-sum + Maximin operator* are widely used in existing works of RARL. We also include a non-robust training baseline (*No-Adv*) to highlight the difference between robust and non-robust training.

Without specific mention, the policies are parametrized by MLPs with two hidden layers. All the agents are trained using policy gradient algorithms with Adam optimizer and the same learning rate. Each plot is computed with 5 policies generated from different random seeds. The episodic reward is evaluated over 48 episodes for each policy. More experiment details are included in Appendix C.

5.2 Highway Merging Task

In highway merging tasks [Leurent, 2018], the protagonist aims to control the ego vehicle (green) to merge into the main lane while avoiding collision with the other yellow vehicle or hitting the end of the ramp. At every timestep, the adversary controls the aggressiveness of the yellow vehicle whose acceleration is proportional to the aggressiveness. The yellow vehicle can only drive in the middle lane, while the ego vehicle can switch lanes.

In this experiment, we aim to answer **(Q1)** and **(Q2)**. We compare our method *RRL-Stack + Stack-PG* with $\alpha = 0.5$, against the benchmark method *zero-sum + Maximin operator* and non-robust training (*No-Adv*) agents. To evaluate the robustness against different environment parameters, we vary the aggressiveness of the yellow vehicle from 0 to 10 and compare the episodic reward of each method.

To answer **(Q1)** about whether *RRL-Stack + Stack-PG* produces challenging yet solvable environments and allows the protagonist to learn robust policies, we visualize the trajectories of the final policies in Fig. 2. For *No-Adv*, since the yellow car does not collide with the protagonist during the training, the protagonist is unaware of the danger of the yellow car. Therefore, *No-Adv* exhibits poor robustness during testing in unseen environments. For *zero-sum + Maximin operator*, the adversary quickly finds the policy to keep blocking the main lane before the ego vehicle enters the lane, which makes the environment completely unsolvable for the protagonist. In this case, the protagonist cannot learn any robust policies but only hits the end of the ramp.

In contrast, for *RRL-Stack + Stack-PG* with $\alpha = 0.5$, since the adversary is not fully adversarial but maximizing the regret of the protagonist, the resulting environments are challenging but still solvable. The protagonist learns robust policies to switch to the middle lane and immediately switch to the leftmost lane to avoid potential collisions. Therefore, *RRL-Stack + Stack-PG* agents exhibit more robustness to the unseen environments than the baselines.

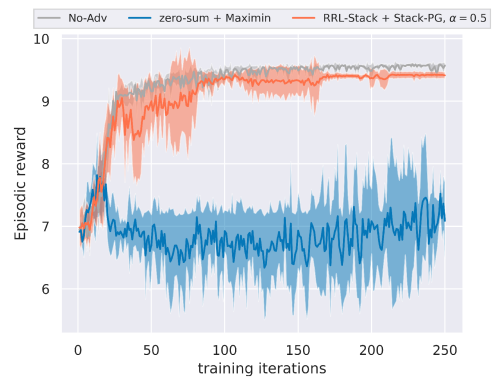


Figure 3: Training curve of highway merging. The x-axis is the training iterations of the protagonist. The y axis is the episodic reward evaluated in the environment without an adversary. The shaded area represents the standard deviation

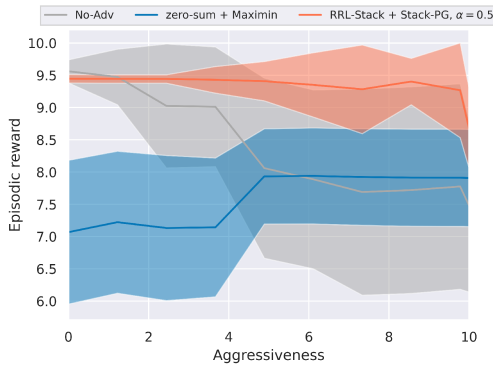


Figure 4: Robustness against different aggressiveness levels. The shaded area represents the standard deviation.

To answer (Q2) about whether our method improves the robustness and training stability, Fig. 3 shows the rewards through the training process. The rewards are evaluated in the same environment without the adversary for a fair comparison. The non-robust training (*No-Adv*) is stable and converges fast because the protagonist is training and evaluating both in the environment without the adversary (however we will observe it is not robust against unseen environments). We observe that after around 20 iterations, the adversary of *zero-sum* + *Maximin operator* quickly learns to make the task unsolvable, so the protagonist’s reward is driven to the lowest possible value for the rest of the training. In contrast, *RRL-Stack* + *Stack-PG* adjusts the difficulty of the environment adaptively to ensure the task remains solvable and the protagonist keeps learning robust policies.

Fig. 4 shows the rewards against different aggressiveness levels. We observe that policies trained with *No-Adv* and *zero-sum* + *Maximin operator* fail to be robust against different aggressiveness levels. In contrast, *RRL-Stack* + *Stack-PG* agents are much more robust against unseen environment parameters during testing. The mean reward of *RRL-Stack* + *Stack-PG* outperforms the baselines by a large margin, and the variance is significantly smaller.

5.3 LunarLander with Actuation Delay

Actuation delay is a common problem in robotic control [Chen *et al.*, 2021]. We modify the LunarLander environment in OpenAI Gym [Brockman *et al.*, 2016] to simulate the effects of actuation delay. The protagonist in LunarLander has 4 discrete

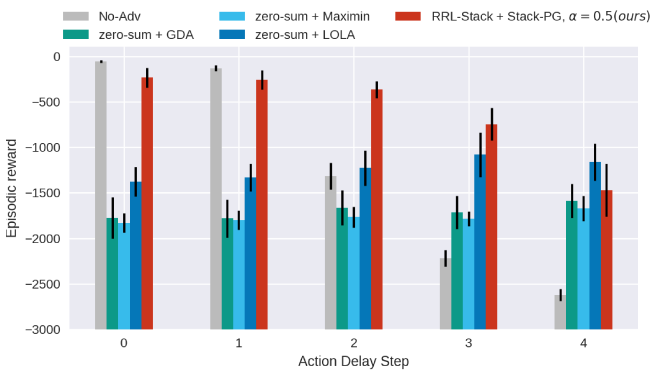


Figure 5: Episodic return with different action delay steps. The black error bar indicates the standard deviation.

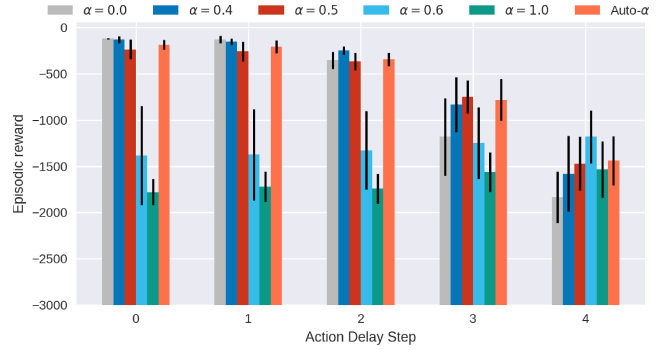


Figure 6: Effects of different α on the episodic rewards.

actions: shut off all engines, turn on the left engine, turn on the right engine and turn on the main engine. The objective is to train a protagonist that is robust to the commanded action being delayed to execute by several steps. During the adversarial training, at each time step, the adversary chooses a number from $\{0, 1, 2, \dots, 10\}$, representing the delay steps of the protagonist’s action. During testing, the delay step is fixed throughout each episode.

In Fig. 5, we show the episodic reward against different action delay steps, compared with the baselines. Only action delay steps from 0 to 4 are shown here since the returns of more delay steps are low for all methods and not meaningful statistically. We find that *RRL-Stack* + *Stack-PG* outperforms the baselines, particularly at delay step 0 to 3.

To answer (Q3) about the effects of α on the robustness of *RRL-Stack* + *Stack-PG*, we study the robustness of $\alpha \in \{0.0, 0.4, 0.5, 0.6, 1.0\}$ as well as our auto-tuning α (*Stack+Auto- α*) in Fig. 6. When $\alpha = 0.4, 0.5$, the protagonist maintains highly robust against different action delay steps, while $\alpha = 0.0$ results in non-robust agents and $\alpha = 0.6, \alpha = 1.0$ produce overly-conservative agents. With *Stack+Auto- α* , the protagonist achieves comparable performance to the best performing $\alpha = 0.4, 0.5$ without fine-tuning.

To study whether *Stack-PG* helps stabilize the training, we use the same *RRL-Stack* game formulation with $\alpha = 0.5$ but apply different learning algorithms in Fig. 7. We test *GDA*, *Maximin* and *Stack-PG* and observe that *Stack-PG* not only stabilizes the training process but also reduces the variance of performance significantly compared with *GDA* and *Maximin*, which again answers (Q2). It is consistent with recent works that have shown that opponent-aware modeling improves the training process stability in Generative Adversarial Networks [Schäfer *et al.*, 2019].

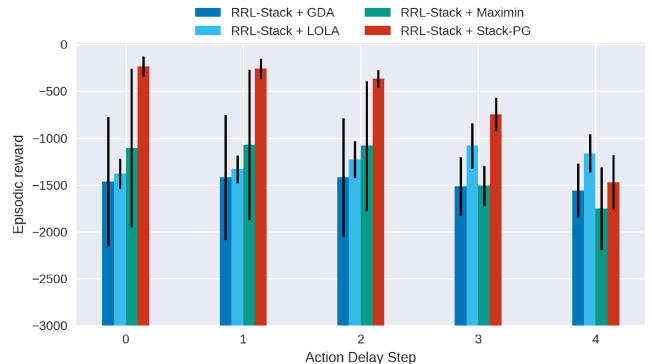


Figure 7: RRL-Stack formulation with different learning algorithms

6 Conclusion

In this work, we study robust reinforcement learning via adversarial training problems. To the best of our knowledge, this is the first work to formalize the sequential nature of deployments of robust RL agents using the Stackelberg game-theoretical formulation. We enable the agent to learn robust policies in progressively challenging environments with the adaptively-regularized adversary. We develop a variant of policy gradient algorithms based on the Stackelberg learning dynamics. In our experiments, we evaluate the robustness of our algorithm on two tasks and demonstrate that our algorithm clearly outperforms the robust and non-robust baselines in single-agent and multi-agent tasks.

Ethical Statement

There are no ethical issues.

Acknowledgments

We gratefully acknowledge support from the National Science Foundation under grants IIS-1849304, CAREER CNS-2047454, and CAREER IIS-2046640.

References

- [Abraham *et al.*, 2012] Ralph Abraham, Jerrold E Marsden, and Tudor Ratiu. *Manifolds, tensor analysis, and applications*, volume 75. Springer Science & Business Media, 2012.
- [Akkaya *et al.*, 2019] Ilge Akkaya, Marcin Andrychowicz, Maciek Chociej, Mateusz Litwin, Bob McGrew, Arthur Petron, Alex Paino, Matthias Plappert, Glenn Powell, Raphael Ribas, et al. Solving rubik’s cube with a robot hand. *arXiv preprint arXiv:1910.07113*, 2019.
- [Başar and Olsder, 1998] Tamer Başar and Geert Jan Olsder. *Dynamic noncooperative game theory*. SIAM, 1998.
- [Brockman *et al.*, 2016] Greg Brockman, Vicki Cheung, Ludwig Pettersson, Jonas Schneider, John Schulman, Jie Tang, and Wojciech Zaremba. Openai gym. *arXiv preprint arXiv:1606.01540*, 2016.
- [Chen *et al.*, 2021] Baiming Chen, Mengdi Xu, Liang Li, and Ding Zhao. Delay-aware model-based reinforcement learning for continuous control. *Neurocomputing*, 450:119–128, 2021.
- [Dennis *et al.*, 2020] Michael Dennis, Natasha Jaques, Eugene Vinitsky, Alexandre Bayen, Stuart Russell, Andrew Critch, and Sergey Levine. Emergent complexity and zero-shot transfer via unsupervised environment design. *arXiv preprint arXiv:2012.02096*, 2020.
- [Désidéri, 2012] Jean-Antoine Désidéri. Multiple-gradient descent algorithm (mgda) for multiobjective optimization. *Comptes Rendus Mathématique*, 350(5-6):313–318, 2012.
- [Ding *et al.*, 2021] Wenhao Ding, Baiming Chen, Bo Li, Kim Ji Eun, and Ding Zhao. Multimodal safety-critical scenarios generation for decision-making algorithms evaluation. *IEEE Robotics and Automation Letters*, 6(2):1551–1558, 2021.
- [Fiez *et al.*, 2020] Tanner Fiez, Benjamin Chasnov, and Lillian Ratliff. Implicit learning dynamics in stackelberg games: Equilibria characterization, convergence analysis, and empirical study. In *International Conference on Machine Learning*, pages 3133–3144. PMLR, 2020.
- [Foerster *et al.*, 2017] Jakob N Foerster, Richard Y Chen, Maruan Al-Shedivat, Shimon Whiteson, Pieter Abbeel, and Igor Mordatch. Learning with opponent-learning awareness. *arXiv preprint arXiv:1709.04326*, 2017.
- [Jin *et al.*, 2020] Chi Jin, Praneeth Netrapalli, and Michael Jordan. What is local optimality in nonconvex-nonconcave minimax optimization? In *International Conference on Machine Learning*, pages 4880–4889. PMLR, 2020.
- [Kamalaruban *et al.*, 2020] Parameswaran Kamalaruban, Yu-Ting Huang, Ya-Ping Hsieh, Paul Rolland, Cheng Shi, and Volkan Cevher. Robust reinforcement learning via adversarial training with langevin dynamics. *arXiv preprint arXiv:2002.06063*, 2020.
- [Leurent, 2018] Edouard Leurent. An environment for autonomous driving decision-making. <https://github.com/eleurent/highway-env>, 2018. Accessed: 2022-05-21.
- [Li *et al.*, 2019] Bin Li, Peng Luo, and Dewen Xiong. Equilibrium strategies for alpha-maxmin expected utility maximization. *SIAM Journal on Financial Mathematics*, 10(2):394–429, 2019.
- [Martens *et al.*, 2012] James Martens, Ilya Sutskever, and Kevin Swersky. Estimating the hessian by back-propagating curvature. *arXiv preprint arXiv:1206.6464*, 2012.
- [Morimoto and Doya, 2005] Jun Morimoto and Kenji Doya. Robust reinforcement learning. *Neural computation*, 17(2):335–359, 2005.
- [Pinto *et al.*, 2017] Lerrel Pinto, James Davidson, Rahul Sukthankar, and Abhinav Gupta. Robust adversarial reinforcement learning. In *International Conference on Machine Learning*, pages 2817–2826. PMLR, 2017.
- [Saad and Schultz, 1986] Youcef Saad and Martin H Schultz. Gmres: A generalized minimal residual algorithm for solving nonsymmetric linear systems. *SIAM Journal on scientific and statistical computing*, 7(3):856–869, 1986.
- [Schäfer *et al.*, 2019] Florian Schäfer, Hongkai Zheng, and Anima Anandkumar. Implicit competitive regularization in gans. *arXiv preprint arXiv:1910.05852*, 2019.
- [Shapley, 1953] Lloyd S Shapley. Stochastic games. *Proceedings of the national academy of sciences*, 39(10):1095–1100, 1953.
- [Shen *et al.*, 2020] Qianli Shen, Yan Li, Haoming Jiang, Zhaoran Wang, and Tuo Zhao. Deep reinforcement learning with robust and smooth policy. In *International Conference on Machine Learning*, pages 8707–8718. PMLR, 2020.
- [Shewchuk and others, 1994] Jonathan Richard Shewchuk et al. An introduction to the conjugate gradient method without the agonizing pain, 1994.

- [Tessler *et al.*, 2019] Chen Tessler, Yonathan Efroni, and Shie Mannor. Action robust reinforcement learning and applications in continuous control. In *International Conference on Machine Learning*, pages 6215–6224. PMLR, 2019.
- [Xu *et al.*, 2020] Mengdi Xu, Wenhao Ding, Jiacheng Zhu, Zuxin Liu, Baiming Chen, and Ding Zhao. Task-agnostic online reinforcement learning with an infinite mixture of gaussian processes. *Advances in Neural Information Processing Systems*, 33:6429–6440, 2020.
- [Xu *et al.*, 2021] Mengdi Xu, Peide Huang, Fengpei Li, Jiacheng Zhu, Xuwei Qi, Kentaro Oguchi, Zhiyuan Huang, Henry Lam, and Ding Zhao. Accelerated policy evaluation: Learning adversarial environments with adaptive importance sampling. *arXiv preprint arXiv:2106.10566*, 2021.
- [Yu *et al.*, 2021] Jing Yu, Clement Gehring, Florian Schäfer, and Animashree Anandkumar. Robust reinforcement learning: A constrained game-theoretic approach. In *Learning for Dynamics and Control*, pages 1242–1254. PMLR, 2021.
- [Zhang *et al.*, 2019] K Zhang, Z Yang, and T Başar. Multi-agent reinforcement learning: A selective overview of theories and algorithms. arxiv e-prints, page. *arXiv preprint arXiv:1911.10635*, 2019.
- [Zhang *et al.*, 2020] Kaiqing Zhang, Bin Hu, and Tamer Basar. On the stability and convergence of robust adversarial reinforcement learning: A case study on linear quadratic systems. *Advances in Neural Information Processing Systems*, 33, 2020.
- [Zhang *et al.*, 2021] Guodong Zhang, Yuanhao Wang, Laurent Lessard, and Roger Grosse. Don’t fix what ain’t broke: Near-optimal local convergence of alternating gradient descent-ascent for minimax optimization. *arXiv preprint arXiv:2102.09468*, 2021.
- [Zhao *et al.*, 2020] Wenshuai Zhao, Jorge Peña Queralta, and Tomi Westerlund. Sim-to-real transfer in deep reinforcement learning for robotics: a survey. In *2020 IEEE Symposium Series on Computational Intelligence (SSCI)*, pages 737–744. IEEE, 2020.