

Leader-Follower Based Online Reinforcement Learning Algorithm in Problem with Hierarchy Decision Makers

Xiaoqian Li, Zonglei Jing, and Peijun Ju

Abstract—This paper presents an online integral reinforcement learning (RL) solution for problems with hierarchy decision makers. Specifically, we reformulate this model as a leader-follower game, in which control input and deterministic disturbance act as decision makers at different levels of hierarchy: The control input plays the role of the leader, while the disturbance plays the role of the follower. The main contributions of this paper can be summarized as follows. First, we introduce online RL to deal with systems that have partially unknown information, meaning that accurate dynamic information is not required. Second, we solve the leader-follower coupled Hamilton-Jacobi (HJ) and Riccati equations approximately online using the derived algorithm. Third, we provide turning laws for cost functions and controllers, which ensure closed-loop stability simultaneously.

Index Terms—Online reinforcement learning algorithm, leader-follower game, hierarchy decision makers

I. INTRODUCTION

Games have been widely used in complex systems involving multiple decision-making players [1]. The behaviors of these players depend on the information that each has compiled about the others, and each seeks to optimize its individual cost. Depending on their roles in practical applications, the involved players may have cooperative or non-cooperative objectives in scenarios that include economics [2, 3] and engineering [4, 5]. Multi-player games can be classified into Nash games or Stackelberg games based on the number of decision-making hierarchies involved. In Nash games, different players select decisions simultaneously without knowing those made by others [6]. Conversely, Stackelberg games involve multi-level decision making, with lower hierarchy players constrained by higher hierarchy players. The mentioned players in Stackelberg game are called the leaders and the followers [7]. Leaders, who are in the higher hierarchy, have complete information on decisions of others and make decisions first. Followers, who are in the lower hierarchy, react to the decisions of leaders. Control

methods for both Nash and Stackelberg games have been applied to dynamical systems to study coordination, conflict, and control among different players [8]. The Stackelberg game has been used in various fields, including multi-input dynamic systems and network communication [9, 10]. In Refs. [11, 12], Stackelberg games with open-loop strategies were introduced to handle multi-layered optimization problems with one leader and one follower. These problems are typically formulated as mixed H_2/H_∞ control, where incomplete system information is considered for the reason of the complex uncertain environments.

Various approaches have been employed to address the Stackelberg game. Xu et al. [13] presented sufficient conditions and necessary conditions for leader-follower games. An algorithm named off-policy reinforcement learning (RL) was applied in Ref. [14] to derive solutions for nonzero-sum games with multi-player. The approach established in Ref. [14] is offline, which places greater demands on system dynamics. However, due to the high uncertainty of the environment, obtaining accurate information on system dynamics poses a challenge. Moreover, players under this control policy cannot adjust their aims in real time. These issues highlight the challenges involved in studying the Stackelberg game.

Reinforcement learning is a promising method for studying multi-player games [15–17], allowing agents to gradually improve their performances by learning new behaviors online with little information about other players. Researchers in control systems have been motivated by the advantages of reinforcement learning and have applied it to find optimal solutions for systems in differential or difference norms. In Ref. [18], online policy was introduced to handle optimization problem with unknown information. Adaptive feedback form controllers were derived by using natural decision methods in Refs. [19, 20]. An optimal control policy for constrained input systems with unknown information was established using policy iteration methods in Ref. [21]. Integral RL was also applied to address tracking problems with constrained inputs [22].

Motivated by previous research on the advantages and outcomes of reinforcement learning, we introduce an online integral RL method to find optimal solutions for mixed H_2/H_∞ problems. In this type of control problem, both control input and deterministic disturbance shape the Stackelberg game, with the control input serving as leader and the disturbance as follower. This paper presents several contributions. First, online RL is introduced to address systems with partially unknown information, eliminating the

Manuscript received: 6 November 2023; revised: 21 December 2023; accepted: 9 January 2024. (Corresponding author: Xiaoqian Li.)

Citation: X. Li, Z. Jing, and P. Ju, Leader-follower based online reinforcement learning algorithm in problem with hierarchy decision makers, *IJICS*, 2024, 29(1), 48–58.

Xiaoqian Li and Peijun Ju are with School of Mathematics and Statistics, Taishan University, Taian 271000, China (e-mail: xiaoqianli0324@126.com; jipj615@163.com).

Zonglei Jing is with State Key Laboratory of Software Development Environment, Beihang University, Beijing 100191, China (e-mail: jingzonglei@163.com).

Digital Object Identifier 10.62678/IJICS202403.10114

need for accurate dynamic information. Second, the leader-follower coupled Hamilton-Jacobi (HJ) and Riccati equations are solved approximately using our derived algorithm. Third, turning laws for cost functions and controllers are provided, ensuring the closed-loop stability.

The rest of this paper is organized as follows. In Section II, we formulate the Stackelberg game for mixed H_2/H_∞ control problems. Section III presents policy iteration methods for Stackelberg games, and Section IV provides value function approximations for these games. We then introduce an online solution for Stackelberg games using RL in Section V. In Section VI, the simulation results verify the validity and feasibility of the proposed method. Finally, we conclude this paper in Section VII.

II. PROBLEM STATEMENT

In this section, we discuss a two-player differential game.

$$\dot{x}(t) = Ax(t) + B_1u(t) + B_2w(t) \quad (1)$$

where $x(t) \in \mathbb{R}^n$ is the state and $x(t_0) = x_0$, t_0 is the initial time and x_0 is the state at initial time; $u(t) \in \mathbb{R}^s$ is the control input; $w(t) \in \mathbb{R}^m$ is the deterministic noise; and t indicates time. $\mathbb{R}^n$ is the family of n -dimensional column vectors, which is the same as $\mathbb{R}^s$ and $\mathbb{R}^m$, n , s , and m are finite positive integers. A , B_1 , and B_2 are constant matrices with appropriate dimensions.

Consider the cost functions

$$J_1(x_0, u) = \int_0^\infty (x^T Q x + u^T R u) dt = \int_0^\infty r_1(x(t), u(t)) dt \quad (2)$$

$$J_2(x_0, u, w) = \int_0^\infty (x^T Q x + u^T R u - \gamma^2 \|w\|^2) dt = \int_0^\infty r_2(x(t), u(t), w(t)) dt \quad (3)$$

where Q is the weighting matrix of state variables, $Q \geq 0$, and R is the weighting matrix of control variables, $R = R^T > 0$. γ is the perturbation factor and γ^* means perturbation factor when the optimal control law $u(t)$ is invoked, and satisfy $\gamma \geq \gamma^* \geq 0$. $\|w\|^2$ means H_2 -norm. For the sake of simplicity, x is equal to $x(t)$, u is equal to $u(t)$, similar treatment is provided elsewhere below.

Definition 1 The existence of the adversarial input $w(t)$ is measured by the perturbed operator L with its H_∞ norm defined by

$$\|L\|_{H_\infty} := \sup_{w \in \mathbb{R}^m} \frac{\left[\int_0^\infty (x^T Q x + u^T R u) dt \right]^{1/2}}{\left[\int_0^\infty (w^T w) dt \right]^{1/2}} < \gamma \quad (4)$$

In this paper, we define the mixed H_2/H_∞ control problem in infinite horizon, where the designed control input $u(t)$ minimizes its corresponding cost function in Eq. (2) under the H_∞ -constraint in Eq. (3). In some practical applications, the adversarial input $w(t)$ maximizes energy gain.

We also provide a general definition of Stackelberg games.

Definition 2 For a Stackelberg game, the sets of admissible strategies for leader and follower are $\mathcal{U}$ and $\mathcal{W}$, respectively. The objective of leader is to minimize J_1 , while follower aims to maximize J_2 .

(1) For a fixed policy $u \in \mathcal{U}$ and a given initial state

$x_0 \in \mathbb{R}^n$, the follower chooses an optimal strategy $w^* \in \mathcal{W}$ such that $J_2(x_0, u, w^*) \geq J_2(x_0, u, w)$.

(2) Assuming an optimal controller w^* exists, which generally depends on the choice of controller u and the initial state x_0 , the leader will select the optimal controller $u \in \mathcal{U}$ to minimize the cost function $J_1(x_0, u, w^*)$ over all $u \in \mathcal{U}$.

Problem 1 We aim to find an optimal feedback strategy pair (u^*, w^*) associated with Eq. (1), for that the following relationships with admissible control pair (u, w) are satisfied.

$$J_1(u^*, w^*) \leq J_1(u, w^*) \quad (5)$$

$$J_2(u, w^*) \geq J_2(u, w) \quad (6)$$

Remark 1 Formula (6) implies that the follower requests the leader to specify its choice based on the follower's optimal policy. In other words, the leader selects a reaction that is identified as the best response. The leader then chooses an input policy from the admissible policies that minimizes its function index, given the follower's designed policy.

III. POLICY ITERATION FOR STACKELBERG GAME

Different hierarchical levels exist in the system due to the denotations of leader and follower. The discussion will begin with the follower.

The optimal cost function of the follower is defined as

$$\begin{aligned} V_2^*(x_0, u, w) &= \max_w J_2(x_0, u, w) = \\ &= \max_w \int_0^\infty (x^T Q x + u^T R u - \gamma^2 \|w\|^2) d\tau = \\ &= \max_w \int_0^\infty r_2(x(t), u(t), w(t)) d\tau \end{aligned} \quad (7)$$

The Hamiltonian of the follower is

$$\begin{aligned} H_2 &= x^T Q x + u^T R u - \gamma^2 \|w\|^2 + \\ &= \frac{\partial V_2}{\partial x}^T (Ax + B_1 u + B_2 w) \end{aligned} \quad (8)$$

where $\frac{\partial V_2}{\partial x}$ denotes the partial derivative of the value function V_2 with respect to x . To enhance visualization, we introduce the alternative operator notation $\nabla V_i \in \mathbb{R}^n$, $i = 1, 2$ to replace $\frac{\partial V_i}{\partial x}$, $i = 1, 2$.

The optimal controller and costate state λ_2 satisfy the following relations

$$w^* = \frac{1}{2}(\gamma^2)^{-1} B_2^T \nabla V_2 \quad (9)$$

$$\dot{\lambda}_2 = -A^T \nabla V_2 - 2Qx \quad (10)$$

The focus now shifts to the leaders. The leader's objective is to minimize Eq. (2) subject to the constraints in Eqs. (1) and (10), which means that the leader should consider the costate of the follower. Then, we have

$$\begin{aligned} V_1^*(x, u, w^*) &= \min_u \int_0^\infty (x^T Q x + u^T R u + y^T \lambda_2) d\tau = \\ &= \min_u \int_0^\infty (r_1(x(t), u(t)) + y^T \lambda_2) d\tau \end{aligned} \quad (11)$$

where V_1^* means the optimal cost function of the leader and y is the costate.

The Hamiltonian of the leader is

$$H_1 = x^T Qx + u^T R u + (\nabla V_1)^T (Ax + B_1 u + B_2 w) + y^T \lambda_2 \quad (12)$$

Similarly, the equilibrium equation and leader's costate equation are derived by Hamilton-Jacobi equation

$$u^* = -\frac{1}{2} R^{-1} B_1^T \nabla V_1 \quad (13)$$

$$\dot{\lambda}_1 = -A^T \nabla V_1 - 2Qx + 2Qy \quad (14)$$

The following Bellman equations are presented in terms of ∇V_1 and ∇V_2 to provide differential equivalents of the value functions for the leader in Eq. (11) and the follower in Eq. (7), where the optimal control policies are inserted into the Hamiltonian.

$$0 = r_1(x, u^*, w^*) + (\nabla V_1)^T (Ax + B_1 u^* + B_2 w^*) + y^T \lambda_2 \quad (15)$$

$$0 = r_2(x, u^*, w^*) + (\nabla V_2)^T (Ax + B_1 u^* + B_2 w^*) \quad (16)$$

By using the infinite integrals in Eqs. (7) and (11), values evaluated along the system trajectories are solutions to Eqs. (15) and (16). The leader's costate is given by

$$\dot{y} = -(2\gamma^2)^{-1} B_2 B_2^T \nabla V_1 + Ay \quad (17)$$

Substituting Eqs. (9) and (13) into Eqs. (15) and (16) yields

$$\begin{aligned} 0 = & x^T Qx + \frac{1}{4} (\nabla V_1)^T B_1 R^{-1} B_1^T \nabla V_1 - \\ & \frac{1}{4} (\nabla V_2)^T B_2 (\gamma^2)^{-1} B_2^T \nabla V_2 + (\nabla V_2)^T (Ax - \\ & \frac{1}{2} B_1 R^{-1} B_1^T \nabla V_1 + \frac{1}{2\gamma^2} B_2 B_2^T \nabla V_2) \end{aligned} \quad (18)$$

$$\begin{aligned} 0 = & x^T Qx + \frac{1}{4} (\nabla V_1)^T B_1 R^{-1} B_1^T \nabla V_1 + (\nabla V_1)^T (Ax - \\ & \frac{1}{2} B_1 R^{-1} B_1^T \nabla V_1 + \frac{1}{2\gamma^2} B_2 B_2^T \nabla V_2) - \\ & y^T (A^T \nabla V_2 + 2Qx) \end{aligned} \quad (19)$$

The following denotations are introduced to rewrite costate variables

$$\nabla V_1 = P_1 x, \nabla V_2 = P_2 x, y = P_3 x \quad (20)$$

where P_1 , P_2 , and P_3 are positive definite symmetric matrices with compatible dimensions. Equations (17)–(19) are rewritten as the following forms

$$\begin{aligned} 0 = & Q + \frac{1}{4} P_1^T B_1 R^{-1} B_1^T P_1 - \frac{1}{4} P_2^T B_2 (\gamma^2)^{-1} B_2^T P_2 + \\ & P_2^T (A - \frac{1}{2} B_1 R^{-1} B_1^T P_1 + \frac{1}{2\gamma^2} B_2 B_2^T P_2) \end{aligned} \quad (21)$$

$$\begin{aligned} 0 = & Q + \frac{1}{4} P_1^T B_1 R^{-1} B_1^T P_1 + P_1^T (A - \frac{1}{2} B_1 R^{-1} B_1^T P_1 + \\ & \frac{1}{2\gamma^2} B_2 B_2^T P_2) - P_3^T (A^T P_2 + 2Q) \end{aligned} \quad (22)$$

$$\begin{aligned} 0 = & -P_3^T (A - \frac{1}{2} B_1 R^{-1} B_1^T P_1 + \frac{1}{2\gamma^2} B_2 B_2^T P_2) - \\ & \frac{1}{2\gamma^2} B_2 B_2^T P_1 + A P_3 \end{aligned} \quad (23)$$

However, solving the Hamilton-Jacobi equations shown in Eqs. (18) and (19) for the leader and follower can be a

significant challenge, and in some cases, obtaining solutions may be impossible. To address this issue, a policy iteration algorithm is introduced to approximate the HJ solution by using iteration algorithm on Bellman Eqs. (15) and (16).

The strategy iteration steps of the Stackelberg game are as follows:

(i) Give initial policies $u^0(x)$ and $w^0(x)$.

(ii) Given $u^k(x)$ and $w^k(x)$, where k represents the k -th step after discretization means, solve for the costs $V_1^k(x)$ and $V_2^k(x)$ by using

$$\begin{aligned} 0 = & r_1(x, u^k(x)) + (\nabla V_1^k)^T (Ax + B_1 u^k + \\ & \frac{1}{2\gamma^2} B_2 B_2^T \nabla V_2^k) + y^T (-A^T \nabla V_2^k - 2Qx) \end{aligned} \quad (24)$$

$$\begin{aligned} 0 = & r_2(x, u^k(x), w^k(x)) + \\ & (\nabla V_2)^T (Ax + B_1 u^k + B_2 w^k) \end{aligned} \quad (25)$$

(iii) Update the control policies using

$$u^{k+1}(x) = -\frac{1}{2} R^{-1} B_1^T \nabla V_1^k \quad (26)$$

$$w^{k+1}(x) = \frac{1}{2\gamma^2} B_2^T \nabla V_2^k \quad (27)$$

Remark 2 The policy iteration algorithm is an offline algorithm that requires a significant amount of dynamic information. This limitation motivates the development of a new online algorithm that can update policies and value functions online while requiring less information about the system dynamics.

IV. VALUE FUNCTION APPROXIMATION FOR STACKELBERG GAME

Based on previous consideration, when the dynamics are unknown, the integral reinforcement learning Bellman equation is presented as follows

$$V_1(x(t-T)) = \int_{t-T}^t r_1(x(t), u(t)) dt + V_1(x(t)) \quad (28)$$

$$V_2(x(t-T)) = \int_{t-T}^t r_2(x(t), u(t), w(t)) dt + V_2(x(t)) \quad (29)$$

where $t \geq 0$ and the time interval $T > 0$.

It should be noted that $V_1^k(x(t))$ can be obtained with the solution of Eq. (24) and $V_2^k(x(t))$ in Eq. (25) can be derived incidentally. Let $V_1'(x)$ and $V_2'(x)$ be admissible feedback control policies of $V_1^k(x(t))$ and $V_2^k(x(t))$, and assume that the value functions are smooth and form a compact set. Thus, $\forall x \in \mathcal{Q}$, where set $\mathcal{Q}$ is included in $\mathbb{R}^n$, $V_1'(x) \geq 0$ and $V_2'(x) \geq 0$. Constant matrix weights W_1 and W_2 are introduced here to approximate the value $V_i(x)$.

$$V_1'(x) = W_1^T \phi_1(x) + \varepsilon_1(x) \quad (30)$$

$$V_2'(x) = W_2^T \phi_2(x) + \varepsilon_2(x) \quad (31)$$

where $\phi_i(x) : \mathbb{R}^n \rightarrow \mathbb{R}^N$, $i = 1, 2$ is one vector of the activation function, the number of hidden layers is N , and $\varepsilon_i(x)$, $i = 1, 2$ is the approximation error, respectively.

Activation functions are chosen to provide a complete and independent basis set, allowing for the uniform approximation of both the value functions and their derivatives.

$$\frac{\partial V'_i}{\partial x} = \nabla V'_i = \left(\frac{\partial \phi_i(x)}{\partial x} \right)^T W_i + \frac{\partial \varepsilon_i}{\partial x} = \nabla \phi_i^T W_i + \nabla \varepsilon_i, \quad i = 1, 2 \quad (32)$$

As demonstrated by the higher-order approximation theorem, with the number of hidden layer neurons N increasing (e.g., $N \rightarrow +\infty$), approximation error $\varepsilon(x)$ and its gradient $\nabla \varepsilon$ both approach zero. Additionally, if the number of layers is fixed, then ε and $\nabla \varepsilon$ are bounded on a compact set with certain constants.

Considering the value function approximation, the Hamiltonians in Eqs. (8) and (12) become

$$\varepsilon_{H_2} = x^T Qx + u^T Ru - \gamma^2 \|w\|^2 + W_2^T \nabla \phi_2(x)(Ax + B_1 u + B_2 w) \quad (33)$$

$$\varepsilon_{H_1} = x^T Qx + u^T Ru + W_1^T \nabla \phi_1(x) \times (Ax + B_1 u + B_2 w) + y^T \lambda_2 \quad (34)$$

where ε_{H_2} and ε_{H_1} are the residual errors for the follower and the leader, respectively.

Under fixed policies u and w , we have

$$\varepsilon_{H_i} = -(\nabla \varepsilon_i)^T (Ax + B_1 u + B_2 w), \quad i = 1, 2 \quad (35)$$

Substituting Eq. (32) into Eqs. (18) and (19), one obtains

$$\begin{aligned} & x^T Qx + \frac{1}{4} W_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) W_1 + \\ & \frac{1}{4} W_2^T \nabla \phi_2(x) B_2 (\gamma^2)^{-1} B_2^T \nabla \phi_2^T(x) W_2 + W_2^T \nabla \phi_2(x) \times \\ & (Ax - \frac{1}{2} B_1 R^{-1} B_1^T \nabla \phi_1^T(x) W_1) = \varepsilon_{H_2} \end{aligned} \quad (36)$$

$$\begin{aligned} & x^T Qx - \frac{1}{4} W_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) W_1 + \\ & W_1^T \nabla \phi_1(x) (Ax + \frac{1}{2\gamma^2} B_2 B_2^T \nabla \phi_2^T(x) W_2) - \\ & y^T (A^T \nabla \phi_2^T(x) W_2 + 2Qx) = \varepsilon_{H_1} \end{aligned} \quad (37)$$

where subscripts H_2 and H_1 mean the second and the first Hamilton-Jacobi equations, respectively. And we have

$$\begin{aligned} \varepsilon_{H_2} = & -\frac{1}{2} W_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \varepsilon_1 - \frac{1}{2\gamma^2} W_2^T \nabla \phi_2(x) \times \\ & B_2 B_2^T \nabla \varepsilon_2 - \frac{1}{4\gamma^2} \nabla \varepsilon_2^T B_2 B_2^T \nabla \varepsilon_2 - \frac{1}{4} \nabla \varepsilon_1^T B_1 \times \\ & R^{-1} B_1^T \nabla \varepsilon_1 - \nabla \varepsilon_2^T Ax + \frac{1}{2} \nabla \varepsilon_2^T B_1 R^{-1} B_1^T \times \\ & \nabla \phi_1^T W_1 + \frac{1}{2} \nabla \varepsilon_2^T B_1 R^{-1} B_1^T \nabla \varepsilon_1 + \\ & \frac{1}{2} W_2^T \nabla \phi_2(x) B_1 R^{-1} B_1^T \nabla \varepsilon_1 \end{aligned} \quad (38)$$

$$\begin{aligned} \varepsilon_{H_1} = & \frac{1}{2} W_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \varepsilon_1 + \frac{1}{4} \nabla \varepsilon_1^T B_1 \times \\ & R^{-1} B_1^T \nabla \varepsilon_1 - \nabla \varepsilon_1^T Ax - \frac{1}{2\gamma^2} \nabla \varepsilon_1^T B_2 B_2^T \times \\ & \nabla \phi_2(x)^T W_2 - \frac{1}{2\gamma^2} \nabla \varepsilon_1^T B_2 B_2^T \nabla \varepsilon_2 - \\ & \frac{1}{2\gamma^2} W_1^T \nabla \phi_1(x) B_2 B_2^T \nabla \varepsilon_2 - y^T A^T \nabla \varepsilon_2 \end{aligned} \quad (39)$$

Lemma 1 For the admissible policy sets $u(x(t))$ and $w(x(t))$, the solution to Eqs. (18) and (19) with least-squares pair exists and changes as N changes. Give the following denotation

$$V_i(x) = W_i^T \phi_i(x), \quad i = 1, 2 \quad (40)$$

Then, as $N \rightarrow +\infty$, we have

- (i) $\sup_{x \in \Omega} |\varepsilon_{H_i}| \rightarrow 0$,
- (ii) $\sup_{x \in \Omega} |\varepsilon_{H_i}| \rightarrow 0$,
- (iii) $\sup_{x \in \Omega} |V_i - V'_i| \rightarrow 0$,
- (iv) $\sup_{x \in \Omega} |\nabla V_i - \nabla V'_i| \rightarrow 0$.

V. ONLINE SOLUTION OF STACKELBERG GAME USING RL

In this section, the issue of developing the approximate optimal adaptive control algorithm is considered, which differs from traditional adaptive control approaches. In the latter, the first step is to design an observer for the cost function and denote it as the value function. The observer obtained is then used to design the feedback strategy. By contrast, the approximate optimal adaptive approach derives the solution of the Stackelberg game online. The outputs of the two players game are as follows

$$\hat{V}_1(x) = \hat{W}_1^T \phi_1(x) \quad (41)$$

$$\hat{V}_2(x) = \hat{W}_2^T \phi_2(x) \quad (42)$$

where $\hat{V}_1$ and $\hat{V}_2$ mean the estimated values of leaders' cost function and followers' cost function, respectively. $\hat{W}_i$ is the current estimated value of the value critic constant matrix weights W_i , $i = 1, 2$.

Considering Eqs. (41) and (42), Hamiltonian equations in Eqs. (8) and (12) become

$$\begin{aligned} H_2(x, \hat{W}_2, u, w) = & x^T Qx + u^T Ru - \gamma^2 \|w\|^2 + \\ & \hat{W}_2^T \nabla \phi_2(x) (Ax + B_1 u + B_2 w) = e_2 \end{aligned} \quad (43)$$

$$\begin{aligned} H_1(x, \hat{W}_1, u, w) = & x^T Qx + u^T Ru + \hat{W}_1^T \nabla \phi_1(x) (Ax + B_1 u + \\ & B_2 w) + y^T (-A^T \nabla \phi_2^T(x) \hat{W}_2 - 2Qx) = e_1 \end{aligned} \quad (44)$$

where e_2 and e_1 represent the estimated values of ε_{H_2} and ε_{H_1} , respectively.

Define the weight estimation error $\tilde{W}_i = W_i - \hat{W}_i$, $i = 1, 2$, then we get

$$\begin{aligned} e_1 = & -\tilde{W}_1^T \nabla \phi_1(x) (Ax + B_1 u + B_2 w) + \varepsilon_{H_1}, \\ e_2 = & -\tilde{W}_2^T \nabla \phi_2(x) (Ax + B_1 u + B_2 w) + \varepsilon_{H_2} \end{aligned} \quad (45)$$

For admissible control policy pair, the selection of $\hat{W}_1$ and $\hat{W}_2$, which aims to minimize the square residual error E , is required.

$$E = \frac{1}{2} \sum_{i=1}^2 e_i^T e_i, \quad i = 1, 2 \quad (46)$$

By calculating, the weights of both follower and leader are selected, and the tuning laws of them are given in normalized gradient descent algorithm

$$\begin{aligned} \dot{\hat{W}}_2 &= -\alpha_2 \frac{\partial E}{\partial \hat{W}_2} = \\ &= -\alpha_2 \frac{\sigma_w}{(1 + \sigma_w^T \sigma_w)^2} [\sigma_w^T \hat{W}_2 + x^T Q x + \\ &+ u^T R u - \gamma^2 \|w\|^2] = \\ &= -\alpha_2 \frac{\sigma_w}{(1 + \sigma_w^T \sigma_w)^2} e_2 \quad (47) \\ \dot{\hat{W}}_1 &= -\alpha_1 \frac{\partial E}{\partial \hat{W}_1} = \\ &= -\alpha_1 \frac{\sigma_u}{(1 + \sigma_u^T \sigma_u)^2} [\sigma_u^T \hat{W}_1 + x^T Q x + u^T R u + \\ &+ y^T (-A^T \nabla \phi_2^T(x) \hat{W}_2 - 2Qx)] = \\ &= -\alpha_1 \frac{\sigma_u}{(1 + \sigma_u^T \sigma_u)^2} e_1 \quad (48) \end{aligned}$$

where $\hat{W}_2$ and $\hat{W}_1$ are the estimated values of W_2 and W_1 , α_2 and α_1 are Levenberg-Marquardt algorithm coefficients, $\sigma_u = \nabla \phi_1(x)(Ax + B_1 u + B_2 w)$, and $\sigma_w = \nabla \phi_2(x)(Ax + B_1 u + B_2 w)$. The previous algorithm is modified for the reason that $(1 + \sigma_w^T \sigma_w)^2$ and $(1 + \sigma_u^T \sigma_u)^2$ are introduced in place of $1 + \sigma_w^T \sigma_w$ and $1 + \sigma_u^T \sigma_u$, respectively. The modified algorithm guarantees the bounded $\frac{\sigma_w}{1 + \sigma_w^T \sigma_w}$ and $\frac{\sigma_u}{1 + \sigma_u^T \sigma_u}$ appearing in Eqs. (47) and (48), respectively.

Equations (33) and (34) can be rewritten as follows

$$x^T Q x + u^T R u - \gamma^2 \|w\|^2 = -W_2^T \nabla \phi_2(x)(Ax + B_1 u + B_2 w) + \varepsilon_{H_2} \quad (49)$$

$$x^T Q x + u^T R u = -W_1^T \nabla \phi_1(x)(Ax + B_1 u + B_2 w) - y^T \lambda_2 + \varepsilon_{H_1} \quad (50)$$

Substituting Eqs. (49) and (50) into Eqs. (47) and (48), respectively, one has

$$\dot{\hat{W}}_2 = -\alpha_2 \frac{\sigma_w}{(1 + \sigma_w^T \sigma_w)^2} [\sigma_w^T (\hat{W}_2 - W_2) + \varepsilon_{H_2}] \quad (51)$$

$$\begin{aligned} \dot{\hat{W}}_1 &= -\alpha_1 \frac{\sigma_u}{(1 + \sigma_u^T \sigma_u)^2} [\sigma_u^T (\hat{W}_1 - W_1) + \varepsilon_{H_1} - \\ &+ y^T A^T \nabla \phi_2^T(x) (\hat{W}_2 - W_2)] \quad (52) \end{aligned}$$

Recalling the denotation that $\tilde{W}_i = W_i - \hat{W}_i$, $i = 1, 2$, the weight estimation error of the dynamic can be written as

$$\dot{\tilde{W}}_2 = \alpha_2 \frac{\sigma_w}{(1 + \sigma_w^T \sigma_w)^2} [-\sigma_w^T \tilde{W}_2 + \varepsilon_{H_2}] \quad (53)$$

$$\dot{\tilde{W}}_1 = \alpha_1 \frac{\sigma_u}{(1 + \sigma_u^T \sigma_u)^2} [-\sigma_u^T \tilde{W}_1 + \varepsilon_{H_1} + y^T A^T \nabla \phi_2^T(x) \tilde{W}_2] \quad (54)$$

Assumption 1 Persistence of excitation (PE) assumption

[23] Denote that $\bar{\sigma}_w = \frac{\sigma_w}{1 + \sigma_w^T \sigma_w}$ and $\bar{\sigma}_u = \frac{\sigma_u}{1 + \sigma_u^T \sigma_u}$ are exciting over the interval $[t, t+T]$, $T > 0$ in persistent way, for example, constants β_i^w and β_i^u , $i = 1, 2$ exist such that for $\forall t$

$$\beta_1^w I \leq S_w, \quad S_w = \int_t^{t+T} \bar{\sigma}_w \bar{\sigma}_w^T d\tau \leq \beta_2^w I \quad (55)$$

$$\beta_1^u I \leq S_u, \quad S_u = \int_t^{t+T} \bar{\sigma}_u \bar{\sigma}_u^T d\tau \leq \beta_2^u I \quad (56)$$

where I indicates the identity matrix with appropriate dimension.

The PE assumption is necessary when using RL to derive system identification in adaptive control. It plays an important role in identifying the critic parameters required to approximate solutions for $V_1(x)$ and $V_2(x)$.

Lemma 2 Given a set of admissible bounded feedback policies $u(x(t))$ and $w(x(t))$, the critics' tuning laws are provided by Eqs. (47) and (48), assuming that $\bar{\sigma}_w$ and $\bar{\sigma}_u$ are persistently exciting. With bounded rationalities $\varepsilon_{\max_1}$ and $\varepsilon_{\max_2}$, the residual errors in Eqs. (33) and (34) are bounded by $\|\varepsilon_{H_i}\| < \varepsilon_{\max_i}$. Therefore, the critic parameter error converges exponentially.

$$\tilde{W}_2(t) \leq \frac{\sqrt{\beta_2^w T}}{\beta_1^w} \{[1 + 2\delta_2 \beta_2^w \alpha_2] \varepsilon_{\max_2}\} \quad (57)$$

$$\tilde{W}_1(t) \leq \frac{\sqrt{\beta_2^u T}}{\beta_1^u} \{[1 + 2\delta_1 \beta_2^u \alpha_1] \varepsilon_{\max_1}\} \quad (58)$$

where δ_1 and δ_2 are positive constants of the order of 1 [24].

Lemma 3 Based on the definition of $V_i(x)$, $i = 1, 2$ in Eq. (40), the control input and adversarial input in Eqs. (9) and (13), respectively, can be written as

$$w^*(x) = \frac{1}{2\gamma^2} B_2^T \nabla \phi_2^T(x) W_2 \quad (59)$$

$$u^*(x) = -\frac{1}{2} R^{-1} B_1^T \nabla \phi_1^T(x) W_1 \quad (60)$$

Then, with the number in the hidden layer tending to infinity, $\sup_{x \in \mathcal{Q}} \|u' - u\| \rightarrow 0$ and $\sup_{x \in \mathcal{Q}} \|w' - w\| \rightarrow 0$ exist.

The algorithm for updating the desired control policies is shown in Eqs. (59) and (60), with unknown matrices W_1 and W_2 . Therefore, denoting the control policies as $\hat{w}(x)$ and $\hat{u}(x)$, we can write them compactly in structured form as follows

$$\hat{w}(x) = \frac{1}{2\gamma^2} B_2^T \nabla \phi_2^T(x) \hat{W}_w \quad (61)$$

$$\hat{u}(x) = -\frac{1}{2} R^{-1} B_1^T \nabla \phi_1^T(x) \hat{W}_u \quad (62)$$

where $\hat{W}_w$ is the current estimated value of W_2 and $\hat{W}_u$ is the current estimated value of W_1 . Actor estimation errors are defined as $\tilde{W}_w = W_2 - \hat{W}_w$ and $\tilde{W}_u = W_1 - \hat{W}_u$.

Definition 3 For given bounds B and T , if relationship $\|x(t)\| \leq B$ for all $t \geq t_0 + T$ holds for any initial value $x_0 \in \mathbb{R}^n$, then $x(t)$ satisfies the speciality of uniformly ultimately bounded.

Definition 4 Given a compact set $\mathcal{Q} \in \mathbb{R}^n$:

(1) The approximation errors ε_1 and ε_2 and their gradients

$\nabla \varepsilon_1$ and $\nabla \varepsilon_2$ are bounded, and the following relationships hold

$$\|\varepsilon_1\| \leq b_{\varepsilon_1}, \|\nabla \varepsilon_1\| \leq b_{\varepsilon_1, x}, \|\varepsilon_2\| \leq b_{\varepsilon_2}, \|\nabla \varepsilon_2\| \leq b_{\varepsilon_2, x} \quad (63)$$

where b_{ε_1} , $b_{\varepsilon_1, x}$, b_{ε_2} , and $b_{\varepsilon_2, x}$ are corresponding boundaries.

(2) The activation functions $\phi_1(x)$ and $\phi_2(x)$ and their gradients $\nabla \phi_1(x)$ and $\nabla \phi_2(x)$ are bounded, and the following relationships hold

$$\begin{aligned} \|\phi_1(x)\| &\leq b_{\phi_1}, \|\nabla \phi_1(x)\| \leq b_{\phi_1, x}, \\ \|\phi_2(x)\| &\leq b_{\phi_2}, \|\nabla \phi_2(x)\| \leq b_{\phi_2, x} \end{aligned} \quad (64)$$

where b_{ϕ_1} , $b_{\phi_1, x}$, b_{ϕ_2} , and $b_{\phi_2, x}$ are corresponding boundaries.

(3) The critic weights W_1 and W_2 are bounded by known constants $W_{1, \max}$ and $W_{2, \max}$

$$\begin{aligned} \|W_1\| &\leq W_{1, \max}, \\ \|W_2\| &\leq W_{2, \max} \end{aligned} \quad (65)$$

The main results of this paper concern the tuning laws for mixed H_2/H_∞ control based on a real-time reinforcement learning algorithm. These tuning laws ensure convergence of the game.

Theorem 1 Considering the dynamics given in Eq. (1) and the Stackelberg game formulated in Section II, the outputs of the two players game are presented in Eqs. (41) and (42), and the control input and disturbance input are given in Eqs. (61) and (62). Let tuning laws be provided by

$$\begin{aligned} \dot{\hat{W}}_2 &= -\alpha_2 \frac{\sigma_{\hat{W}}}{(1 + \sigma_{\hat{W}}^T \sigma_{\hat{W}})^2} [\sigma_{\hat{W}}^T \hat{W}_2 + x^T Q x + \\ &\quad \hat{u}^T R \hat{u} - \gamma^2 \|\hat{w}\|^2] \end{aligned} \quad (66)$$

$$\begin{aligned} \dot{\hat{W}}_1 &= -\alpha_1 \frac{\sigma_{\hat{u}}}{(1 + \sigma_{\hat{u}}^T \sigma_{\hat{u}})^2} [\sigma_{\hat{u}}^T \hat{W}_1 + x^T Q x + \hat{u}^T R \hat{u} + \\ &\quad y^T (-A^T \nabla \phi_2^T(x) \hat{W}_2 - 2Qx)] \end{aligned} \quad (67)$$

where $\sigma_{\hat{W}} = \nabla \phi_2(x)(Ax + B_1 \hat{u} + B_2 \hat{w})$ and $\sigma_{\hat{u}} = \nabla \phi_1(x)(Ax + B_1 \hat{u} + B_2 \hat{w})$.

Tuning laws $\dot{\hat{W}}_w$ and $\dot{\hat{W}}_u$ for follower's and leader's actor are designed as

$$\dot{\hat{W}}_w = -\alpha_4 \left[(\kappa_4 \hat{W}_w - \kappa_3 \bar{\sigma}_{\hat{W}}^T \hat{W}_2) - \frac{1}{4} \bar{D}_2(x) \hat{W}_w m_w^T \hat{W}_2 \right] \quad (68)$$

$$\begin{aligned} \dot{\hat{W}}_u &= -\alpha_3 \left[(\kappa_2 \hat{W}_u - \kappa_1 \bar{\sigma}_{\hat{u}}^T \hat{W}_1) - \frac{1}{4} (\bar{D}_1(x) \hat{W}_u m_u^T \hat{W}_2 + \right. \\ &\quad \left. \bar{D}_1(x) \hat{W}_u m_u^T \hat{W}_1) + \frac{1}{2\gamma^2} \nabla \phi_1 B_2 B_2^T P_2 y \right] \end{aligned} \quad (69)$$

where $\bar{D}_1(x) = \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x)$, $\bar{D}_2(x) = \nabla \phi_2(x) B_2$
 $(\gamma^2)^{-1} B_2^T \nabla \phi_2^T(x)$, $m_w = \frac{\sigma_{\hat{W}}}{(1 + \sigma_{\hat{W}}^T \sigma_{\hat{W}})^2}$, $m_u = \frac{\sigma_{\hat{u}}}{(1 + \sigma_{\hat{u}}^T \sigma_{\hat{u}})^2}$,
 $\bar{\sigma}_{\hat{W}} = \frac{\sigma_{\hat{W}}}{1 + \sigma_{\hat{W}}^T \sigma_{\hat{W}}}$, $\bar{\sigma}_{\hat{u}} = \frac{\sigma_{\hat{u}}}{1 + \sigma_{\hat{u}}^T \sigma_{\hat{u}}}$, respectively. $\kappa_i > 0$ and $\alpha_i > 0, \forall i = 1, 2, 3, 4$ are tuning parameters, which are detailedly presented in the proof of Theorem 1 in Appendix A. Suppose that $\bar{\sigma}_{\hat{W}}$ and $\bar{\sigma}_{\hat{u}}$ are persistently exciting, a constant number N_0 exists such that, for any $N > N_0$, the following mentioned notions, such as the system state, the critic errors $\hat{W}_i$, $i = 1, 2$, and the actor errors $\hat{W}_w$ and $\hat{W}_u$, are uniformly ultimately bounded.

Theorem 2 If the hypotheses proposed in Theorem 1 are met, the following results hold

(i) $H_1(x, \hat{W}_1, \hat{w}'(x), \hat{u}'(x))$ and $H_2(x, \hat{W}_2, \hat{w}'(x), \hat{u}'(x))$ are uniformly ultimately bounded.

$$\hat{w}'(x) = \frac{1}{2\gamma^2} B_2^T \nabla \phi_2^T(x) \hat{W}_2 \quad (70)$$

$$\hat{u}'(x) = -\frac{1}{2} R^{-1} B_1^T \nabla \phi_1^T(x) \hat{W}_1 \quad (71)$$

where $\hat{w}'$ and $\hat{u}'$ are approximate values of $\hat{w}$ and $\hat{u}$, respectively. $\hat{W}_1$ and $\hat{W}_2$ converge to the approximate coupled HJ solution.

(ii) $\hat{w}(x)$ and $\hat{u}(x)$ in Eqs. (61) and (62) converge to the approximate solution of the Stackelberg game.

The proof of Theorem 2 is presented in Appendix B.

VI. SIMULATION

The 2-degree of freedom (2-DOF) robot serves as a basic control platform for the manipulator. The system is a planar 2-DOF with two drives and three unactuated swivel joints. Two servomotors, mounted at a fixed distance on the rotary motion servo platform, control a four-bar system. The goal of 2-DOF robot trajectory tracking is to regulate the $x-y$ position through a 4-link end-actuator. To achieve this, a controller must be designed to regulate the plane coordinates of the end-actuator and to drive the two arms connected to the motor to realize the motion of the other two arms. The trajectory is represented by the end-effector.

Based on the output angular displacements θ_A and θ_B of the two servo systems detected by the encoder, the corresponding control voltage V_m is calculated at each moment. The power module then controls the servo system. The principle diagram of the control system is shown in Fig. 1.

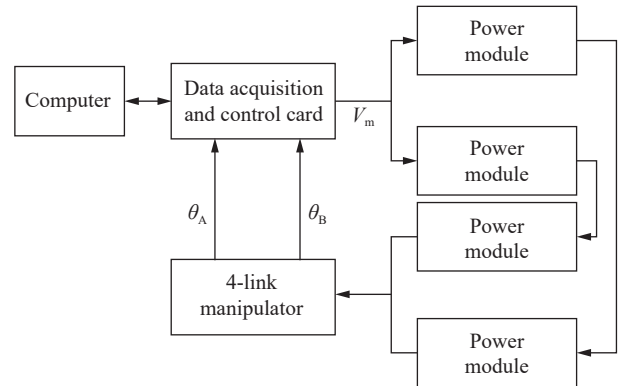


Figure 1 Schematic diagram of 2-DOF robot control system.

Considering the system with the form shown in Eq. (1), we have

$$\begin{aligned} \dot{x}(t) &= \begin{bmatrix} 0.000 & 1.000 \\ -835.210 & -39.882 \end{bmatrix} x(t) + \\ &\quad \begin{bmatrix} 0.000 \\ -835.210 \end{bmatrix} u(t) + \begin{bmatrix} 0.000 \\ 1.000 \end{bmatrix} w(t) \end{aligned} \quad (72)$$

with matrices $Q = R = I$, $\gamma = 0.09$, and tuning gains $\alpha_1 = 5$, $\alpha_2 = 1$, $\alpha_3 = 1$, and $\alpha_4 = 1$.

The computer provides the expected trajectory of the state. Figure 2 compares the state obtained from the online reinforcement learning algorithm with the ideal state. It is clear that the estimated state can quickly follow the ideal, demonstrating the effectiveness of the algorithm discussed in this paper.

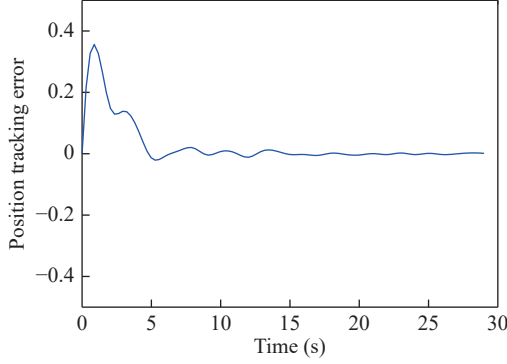


Figure 2 Tracking error.

VII. CONCLUSION

This paper presents an online integral RL algorithm for finding the optimal solution to a problem with hierarchical decision makers. In this model, the leader is the control input and the follower is the disturbance. Both the control input and the deterministic disturbance shape the Stackelberg game. Full dynamic information of the system is not required, as online solution strategies are applied, making explicit solutions unnecessary. The designed approximators for the leader and follower are used to learn decision policies online for the purpose of team decision making. This online learning approach is suitable for varying team dynamics. As shown in this paper, the algorithms presented here offer many advantages in dealing with dynamic systems, particularly mixed H_2/H_∞ problems.

APPENDIX A PROOF OF THEOREM 1

Proof The convergence proof is presented by introducing Lyapunov analysis.

Specifically, we consider the Lyapunov function $L(x)$

$$L(x) = V_1(x) + V_2(x) + \frac{1}{2} \tilde{W}_1^T \alpha_1^{-1} \tilde{W}_1 + \frac{1}{2} \tilde{W}_2^T \alpha_2^{-1} \tilde{W}_2 + \frac{1}{2} \tilde{W}_w^T \alpha_3^{-1} \tilde{W}_w + \frac{1}{2} \tilde{W}_u^T \alpha_4^{-1} \tilde{W}_u \quad (A1)$$

where $V_1(x)$ and $V_2(x)$ are approximate solutions of Eqs. (18) and (19) and are given in Lemma 1.

It is easy to obtain the time derivative as follows

$$\dot{L}(x) = \dot{V}_1(x) + \dot{V}_2(x) + \tilde{W}_1^T \alpha_1^{-1} \dot{\tilde{W}}_1 + \tilde{W}_2^T \alpha_2^{-1} \dot{\tilde{W}}_2 + \tilde{W}_w^T \alpha_3^{-1} \dot{\tilde{W}}_w + \tilde{W}_u^T \alpha_4^{-1} \dot{\tilde{W}}_u \quad (A2)$$

The key step is calculating the terms in $\dot{L}(x)$. For the first term, differentiate Eq. (40) when $i = 1$, then add and subtract $\frac{1}{2} W_1^T \bar{D}_1(x) W_1$ and $\frac{1}{2} W_2^T \bar{E}_2(x) W_2$, we have

$$\begin{aligned} \dot{V}_1(x) &= W_1^T \sigma_1 + \frac{1}{2} W_1^T \bar{D}_1(x) \tilde{W}_u - \frac{1}{2} W_1^T \bar{E}_2(x) \tilde{W}_w + \\ &\quad \dot{\epsilon}_1(x) \end{aligned} \quad (A3)$$

where

$$\begin{aligned} \sigma_1 &= \nabla \phi_1(x)(Ax + B_1 u' + B_2 w'), \\ \bar{E}_2(x) &= \nabla \phi_1(x) B_2 (\gamma^2)^{-1} B_2^T \nabla \phi_1^T(x), \\ \dot{\epsilon}_1(x) &= \nabla \epsilon_1^T(x) [Ax - \frac{1}{2} B_1 R^{-1} B_1^T \nabla \phi_1^T \hat{W}_u + \\ &\quad \frac{1}{2} B_2 (\gamma^2)^{-1} B_2^T \nabla \phi_2^T \hat{W}_w], \\ W_1^T \sigma_1 &= -x^T Qx - \frac{1}{4} W_1^T \bar{D}_1 W_1 + y^T (A^T \nabla \phi_2^T(x) W_2 + \\ &\quad 2Qx) + \epsilon_{HJ1} \end{aligned} \quad (A4)$$

Using similar method to deal with the second term, we have

$$\begin{aligned} \dot{V}_2(x) &= W_2^T \sigma_2 + \frac{1}{2} W_2^T \bar{E}_1(x) \tilde{W}_w + \\ &\quad \frac{1}{2} W_2^T \bar{D}_2(x) \tilde{W}_w + \dot{\epsilon}_2(x) \end{aligned} \quad (A5)$$

where

$$\begin{aligned} \sigma_2 &= \nabla \phi_2(x)(Ax + B_1 u' + B_2 w'), \\ \bar{E}_1(x) &= \nabla \phi_2(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x), \\ \dot{\epsilon}_2(x) &= \nabla \epsilon_2^T(x) [Ax - \frac{1}{2} B_1 R^{-1} B_1^T \nabla \phi_1^T \hat{W}_u + \\ &\quad \frac{1}{2} B_2 (\gamma^2)^{-1} B_2^T \nabla \phi_2^T \hat{W}_w], \\ W_2^T \sigma_2 &= -x^T Qx - \frac{1}{4} W_1^T \bar{D}_1 W_1 + \frac{1}{4} W_2^T \bar{D}_2 W_2 + \epsilon_{HJ2} \end{aligned} \quad (A6)$$

Adding $\dot{V}_1(x)$ and $\dot{V}_2(x)$, thus we get

$$\begin{aligned} \dot{L}_v &= \dot{L}_v + \frac{1}{2} W_1^T \bar{D}_1(x) \tilde{W}_u - \frac{1}{2} W_1^T \bar{E}_2(x) \tilde{W}_w + \\ &\quad \frac{1}{2} W_2^T \bar{E}_2(x) \tilde{W}_w + \frac{1}{2} W_2^T \bar{D}_2(x) \tilde{W}_w + \\ &\quad \dot{\epsilon}_1(x) + \dot{\epsilon}_2(x) \end{aligned} \quad (A7)$$

for the sake of simplicity

$$\begin{aligned} \dot{L}_v &= -x^T Qx - \frac{1}{4} W_1^T \bar{D}_1 W_1 + y^T (A^T \nabla \phi_2^T(x) W_2 + \\ &\quad 2Qx) + \epsilon_{HJ1} - x^T Qx - \frac{1}{4} W_1^T \bar{D}_1 W_1 + \\ &\quad \frac{1}{4} W_2^T \bar{D}_2 W_2 + \epsilon_{HJ2} \end{aligned} \quad (A8)$$

Combining the tuning law in Eq. (67) and the error definitions, we can obtain the third term

$$\begin{aligned} \dot{L}_3 &= \tilde{W}_1^T \alpha_1^{-1} \dot{\tilde{W}}_1 = \\ &\quad \tilde{W}_1^T m_u [\sigma_u^T \hat{W}_1 + x^T Qx + \hat{u}^T R \hat{u} + \\ &\quad y^T (-A^T \phi_2^T(x) \hat{W}_2 - 2Qx)] = \\ &\quad \dot{L}_3 + \tilde{W}_1^T m_u y^T A^T \nabla \phi_2^T(x) \tilde{W}_2 + \frac{1}{4} \tilde{W}_1^T m_u \times \\ &\quad [\tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u + 2W_1^T \bar{E}_2(x) \tilde{W}_w] \end{aligned} \quad (A9)$$

where

$$\dot{\tilde{L}}_3 = \tilde{W}_1^T \bar{\sigma}_{\hat{u}} \left[-\bar{\sigma}_{\hat{u}} \tilde{W}_1 + \frac{\varepsilon_{HJ1}}{1 + \sigma_{\hat{u}}^T \sigma_{\hat{u}}} \right] \quad (A10)$$

By using the similar method, the forth term becomes

$$\begin{aligned} \dot{\tilde{L}}_4 = & \dot{\tilde{L}}_4 + \frac{1}{4} \tilde{W}_2^T m_w [\tilde{W}_w^T \bar{D}_2(x) \tilde{W}_w + \tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u + \\ & 2\tilde{W}_u^T \bar{E}_1(x) W_2 - 2\tilde{W}_u^T \bar{D}_1(x) W_1] \end{aligned} \quad (A11)$$

where

$$\dot{\tilde{L}}_4 = \tilde{W}_2^T \bar{\sigma}_{\hat{w}} \left[-\bar{\sigma}_{\hat{w}} \tilde{W}_2 + \frac{\varepsilon_{HJ2}}{1 + \sigma_{\hat{w}}^T \sigma_{\hat{w}}} \right] \quad (A12)$$

In order to obtain the tuning laws of actors of follower and leader, these two terms $\tilde{W}_u$ and $\tilde{W}_w$ are grouped by multiplying with estimated values. This yields

$$\begin{aligned} \dot{\tilde{L}}(x) = & \dot{\tilde{L}}_v + \dot{\tilde{L}}_3 + \dot{\tilde{L}}_4 + \dot{\varepsilon}_1(x) + \dot{\varepsilon}_2(x) + \frac{1}{2} W_1^T \bar{D}_1(x) \times \\ & \tilde{W}_u - \frac{1}{2} W_1^T \bar{E}_2(x) \tilde{W}_w + \frac{1}{2} W_2^T \bar{E}_2(x) \tilde{W}_w + \\ & \frac{1}{2} W_2^T \bar{D}_2(x) \tilde{W}_w + \tilde{W}_1^T m_u y^T A^T \nabla \phi_2^T(x) \tilde{W}_2 + \\ & \frac{1}{2} W_1^T \bar{E}_2(x) \tilde{W}_w m_u \tilde{W}_1 + \frac{1}{2} W_2^T \bar{E}_1(x) \tilde{W}_u m_w \times \\ & \tilde{W}_2 - \frac{1}{2} W_1^T \bar{D}_1(x) \tilde{W}_u m_w \tilde{W}_2 + \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) \times \\ & W_1 m_u \tilde{W}_1 - \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) W_1 m_u W_1 + \frac{1}{4} \tilde{W}_u^T \times \\ & \bar{D}_1(x) \tilde{W}_u m_u W_1 + \frac{1}{4} \tilde{W}_w^T \bar{D}_2(x) W_2 m_w \tilde{W}_2 - \\ & \frac{1}{4} \tilde{W}_w^T \bar{D}_2(x) W_2 m_w W_2 + \frac{1}{4} \tilde{W}_w^T \bar{D}_2(x) \tilde{W}_w m_w \times \\ & W_2 + \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) W_1 m_w \tilde{W}_2 - \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) \times \\ & W_1 m_w W_2 + \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u m_w W_2 - \tilde{W}_w^T [\alpha_4^{-1} \times \\ & \dot{\tilde{W}}_w - \frac{1}{4} \bar{D}_2(x) \hat{W}_w m_w \hat{W}_2] - \tilde{W}_u^T [\alpha_3^{-1} \dot{\tilde{W}}_u - \\ & \frac{1}{4} \bar{D}_1(x) \hat{W}_u m_u \hat{W}_1 - \frac{1}{4} \bar{D}_1(x) \hat{W}_u m_w \hat{W}_w + \\ & \frac{1}{2\gamma^2} \nabla \phi_1 B_2 B_2^T P_2 y] \end{aligned} \quad (A13)$$

Considering the last two terms, the following tuning laws for the leader and follower can be designed to simplify the equation

$$\dot{\tilde{W}}_w = -\alpha_4 [(\kappa_4 \hat{W}_w - \kappa_3 \bar{\sigma}_{\hat{w}}^T \hat{W}_2) - \frac{1}{4} \bar{D}_2(x) \hat{W}_w m_w^T \hat{W}_2] \quad (A14)$$

$$\begin{aligned} \dot{\tilde{W}}_u = & -\alpha_3 [(\kappa_2 \hat{W}_u - \kappa_1 \bar{\sigma}_{\hat{u}}^T \hat{W}_1) - \frac{1}{4} (\bar{D}_1(x) \hat{W}_u m_u^T \hat{W}_2 + \\ & \bar{D}_1(x) \hat{W}_u m_u^T \hat{W}_1) + \frac{1}{2\gamma^2} \nabla \phi_1 B_2 B_2^T P_2 y] \end{aligned} \quad (A15)$$

Then, the time derivative of Lyapunov function can be rewritten as

$$\dot{\tilde{L}}(x) = -x^T Q x - \frac{1}{4} W_1^T \bar{D}_1 W_1 + y^T (A^T \nabla \phi_2^T(x) W_2 +$$

$$2Qx) + \varepsilon_{HJ1} - x^T Q x - \frac{1}{4} W_1^T \bar{D}_1 W_1 + \frac{1}{4} W_2^T \times$$

$$\bar{D}_2 W_2 + \varepsilon_{HJ2} + \tilde{W}_1^T m_u y^T A^T \nabla \phi_2^T(x) \tilde{W}_2 +$$

$$\tilde{W}_1^T \bar{\sigma}_{\hat{u}} \left[-\bar{\sigma}_{\hat{u}} \tilde{W}_1 + \frac{\varepsilon_{HJ1}}{1 + \sigma_{\hat{u}}^T \sigma_{\hat{u}}} \right] + \dot{\varepsilon}_1(x) \tilde{W}_2^T \times$$

$$\bar{\sigma}_{\hat{w}} \left[-\bar{\sigma}_{\hat{w}} \tilde{W}_2 + \frac{\varepsilon_{HJ2}}{1 + \sigma_{\hat{w}}^T \sigma_{\hat{w}}} \right] + \dot{\varepsilon}_2(x) + \frac{1}{2} W_1^T \times$$

$$\bar{D}_1(x) \tilde{W}_u - \frac{1}{2} W_1^T \bar{E}_2(x) \tilde{W}_w + \frac{1}{2} W_2^T \bar{E}_2(x) \times$$

$$\tilde{W}_w + \frac{1}{2} W_2^T \bar{D}_2(x) \tilde{W}_w + \frac{1}{2} W_1^T \bar{E}_2(x) \tilde{W}_w m_u \tilde{W}_1 +$$

$$\frac{1}{2} W_2^T \bar{E}_1(x) \tilde{W}_u m_w \tilde{W}_2 - \frac{1}{2} W_1^T \bar{D}_1(x) \tilde{W}_u m_w \tilde{W}_2 +$$

$$\frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) W_1 m_u \tilde{W}_1 - \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) W_1 m_u W_1 +$$

$$\frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u m_u W_1 + \frac{1}{4} \tilde{W}_w^T \bar{D}_2(x) W_2 m_w \tilde{W}_2 -$$

$$\frac{1}{4} \tilde{W}_w^T \bar{D}_2(x) W_2 m_w W_2 + \frac{1}{4} \tilde{W}_w^T \bar{D}_2(x) \tilde{W}_w m_w W_2 +$$

$$\frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) W_1 m_w \tilde{W}_2 - \frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) W_1 m_w W_2 +$$

$$\frac{1}{4} \tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u m_w W_2 + \tilde{W}_w^T \kappa_4 W_2 - \tilde{W}_w^T \kappa_4 \tilde{W}_w -$$

$$\tilde{W}_w^T \kappa_3 \bar{\sigma}_{\hat{w}}^T W_2 + \tilde{W}_w^T \kappa_3 \bar{\sigma}_{\hat{w}}^T \tilde{W}_w + \tilde{W}_u^T \kappa_2 W_1 -$$

$$\tilde{W}_u^T \kappa_2 \tilde{W}_u - \tilde{W}_u^T \kappa_1 \bar{\sigma}_{\hat{u}}^T W_1 + \tilde{W}_u^T \kappa_1 \bar{\sigma}_{\hat{u}}^T \tilde{W}_u \quad (A16)$$

Here comes the norm bounds. Based on the definition of uniform ultimate boundedness, Eqs. (A4) and (A6) become

$$\|\dot{\varepsilon}_1(x)\| < b_{\varepsilon_1 x} A \|x\| + \frac{1}{2} b_{\varepsilon_1 x} B_1^2 b_{\phi_1 x} \sigma_{\min}(R) (\|W_1\| +$$

$$\|\tilde{W}_u\|) + \frac{1}{2} b_{\varepsilon_1 x} B_2^2 b_{\phi_2 x} \sigma_{\min}(\gamma^2) (\|W_2\| +$$

$$\|\tilde{W}_w\|) \quad (A17)$$

$$\|\dot{\varepsilon}_2(x)\| < b_{\varepsilon_2 x} A \|x\| + \frac{1}{2} b_{\varepsilon_2 x} B_1^2 b_{\phi_1 x} \sigma_{\min}(R) (\|W_1\| +$$

$$\|\tilde{W}_u\|) + \frac{1}{2} b_{\varepsilon_2 x} B_2^2 b_{\phi_2 x} \sigma_{\min}(\gamma^2) (\|W_2\| +$$

$$\|\tilde{W}_w\|) \quad (A18)$$

where $\sigma_{\min}$ means the inversion of matrix R .

Given $\bar{\varepsilon}_i > 0$, $i = 1, 2$, $K_0(\bar{\varepsilon}_1)$, and $K_0(\bar{\varepsilon}_2)$, where K_0 is the hidden layer unit, there are $\sup_{x \in \Omega} \|\varepsilon_{HJ1}\| < \bar{\varepsilon}_1$ and $\sup_{x \in \Omega} \|\varepsilon_{HJ2}\| < \bar{\varepsilon}_2$. Then assuming $N > K_0$, define

$$\tilde{\Xi} = [x \quad \bar{\sigma}_{\hat{u}}^T \tilde{W}_1 \quad \bar{\sigma}_{\hat{w}}^T \tilde{W}_2 \quad \tilde{W}_u \quad \tilde{W}_w \quad y]^T \quad (A19)$$

there is

$$\begin{aligned}
\dot{L}(x) &= \frac{1}{2}\|W_1\|^2\|\bar{D}_1\| + \frac{1}{4}\|W_2\|^2\|\bar{D}_2\| + \bar{\varepsilon}_1 + \\
&\quad \bar{\varepsilon}_2 \frac{1}{2} b_{\varepsilon_1 x} B_1^2 b_{\phi_1 x} \sigma_{\min}(R) \|W_1\| + \frac{1}{2} b_{\varepsilon_1 x} B_2^2 \times \\
&\quad b_{\phi_2 x} \sigma_{\min}(\gamma^2) \|W_2\| + \frac{1}{2} b_{\varepsilon_2 x} B_1^2 b_{\phi_1 x} \sigma_{\min}(R) \times \\
&\quad \|W_1\| + \frac{1}{2} b_{\varepsilon_2 x} B_2^2 b_{\phi_2 x} \sigma_{\min}(\gamma^2) \|W_2\| + \\
&\quad \tilde{L}^T \begin{bmatrix} -2Q & 0 & 0 & 0 & 0 & 0 \\ 0 & -I & 0 & m_{24} & m_{25} & m_{26} \\ 0 & 0 & -I & m_{34} & m_{35} & 0 \\ 0 & m_{42} & m_{43} & 0 & 0 & 0 \\ 0 & m_{52} & m_{53} & 0 & m_{55} & 0 \\ 0 & m_{62} & 0 & 0 & 0 & 0 \end{bmatrix} \tilde{L} + \\
&\quad \tilde{L}^T \begin{bmatrix} (b_{\varepsilon_1 x} + b_{\varepsilon_2 x})A \\ \frac{\varepsilon_{HJ_1}}{m_{s1}} \\ \frac{\varepsilon_{HJ_2}}{m_{s2}} \\ \bar{m}_4 \\ \bar{m}_5 \\ A^T \nabla \phi_2^T(x) W_2 \end{bmatrix} = \bar{\varepsilon}_1 + \bar{\varepsilon}_2 + \bar{M} + \\
&\quad \tilde{L}^T M \tilde{L} + \tilde{L}^T \tilde{M} \quad (A20)
\end{aligned}$$

where M is a matrix, $\bar{M}$ and $\tilde{M}$ are bounded in $\bar{M}_{\max}$ and $\tilde{M}_{\max}$, and

$$\begin{aligned}
m_{24} &= \frac{1}{2} \kappa_1 + \frac{1}{8m_{s1} \bar{D}_1(x) W_1} = m_{42}^T, \\
m_{25} &= \frac{1}{4m_{s1}} \bar{E}_2(x) W_1 = m_{52}^T, \\
m_{26} &= \frac{1}{2m_{s1}} A^T \nabla \phi_2^T(x) \tilde{W}_2 = m_{62}^T, \\
m_{34} &= \frac{1}{4m_{s2}} \bar{E}_1(x) W_2 - \frac{1}{4m_{s2}} \bar{D}_1(x) W_1 + \\
&\quad \frac{1}{4m_{s2}} \bar{D}_1(x) W_2 = m_{43}^T, \\
m_{35} &= \frac{1}{2} \kappa_3 + \frac{1}{8m_{s1}} \bar{D}_2(x) W_2 = m_{53}^T, \\
m_{44} &= -\kappa_2 + \frac{1}{4} \bar{D}_1(x) m_u W_1 + \frac{1}{4} \bar{D}_1(x) m_w W_2, \\
m_{55} &= -\kappa_4 + \frac{1}{4} \bar{D}_2(x) m_w W_2, \\
m_{s1} &= 1 + \sigma_{\hat{u}}^T \sigma_{\hat{u}}, \\
m_{s2} &= 1 + \sigma_{\hat{w}}^T \sigma_{\hat{w}}, \\
\bar{m}_4 &= \kappa_2 W_1 - \kappa_1 \bar{\sigma}_{\hat{u}}^T W_1 + \frac{1}{2} \bar{D}_1(x) W_1 + \\
&\quad \frac{1}{4} \bar{D}_1(x) W_1 m_u W_1 - \frac{1}{4} \bar{D}_1(x) W_1 m_w W_2, \\
\bar{m}_5 &= \kappa_4 W_2 - \kappa_3 \bar{\sigma}_{\hat{w}}^T W_2 + \frac{1}{2} \bar{E}_2(x) W_1 + \frac{1}{2} \bar{E}_2(x) W_2 + \\
&\quad \frac{1}{2} \bar{D}_2(x) W_2 - \frac{1}{4} \bar{D}_2(x) W_2 m_w W_2, \\
\bar{M} &= \frac{1}{2} \|W_1\|^2 \|\bar{D}_1\| + \frac{1}{4} \|W_2\|^2 \|\bar{D}_2\| + \frac{1}{2} b_{\varepsilon_1 x} B_1^2 b_{\phi_1 x} \times \\
&\quad \sigma_{\min}(R) \|W_1\| + \frac{1}{2} b_{\varepsilon_1 x} B_2^2 b_{\phi_2 x} \sigma_{\min}(\gamma^2) \|W_2\| + \\
&\quad \frac{1}{2} b_{\varepsilon_2 x} B_1^2 b_{\phi_1 x} \sigma_{\min}(R) \|W_1\| + \\
&\quad \frac{1}{2} b_{\varepsilon_2 x} B_2^2 b_{\phi_2 x} \sigma_{\min}(\gamma^2) \|W_2\| \quad (A21)
\end{aligned}$$

According to the bounds given in Definition 4, $\bar{M}$ and $\tilde{M}$ are bounded in $\bar{M}_{\max}$ and $\tilde{M}_{\max}$, respectively.

Now, the matrix M can be compacted and rewritten as

$$M = \begin{bmatrix} -2Q & 0 & 0 & 0 \\ 0 & -I_2 & M'_{23} & M'_{24} \\ 0 & M'_{32} & M'_{33} & 0 \\ 0 & M'_{42} & 0 & 0 \end{bmatrix} = \begin{bmatrix} -2Q & 0 & 0 \\ 0 & -I_2 & M_{23} \\ 0 & M_{32} & M_{33} \end{bmatrix} \quad (A22)$$

where

$$\begin{aligned}
-I_2 &= \begin{bmatrix} -I & 0 \\ 0 & -I \end{bmatrix}, \\
M'_{23} &= \begin{bmatrix} m_{24} & m_{25} \\ m_{34} & m_{35} \end{bmatrix}, \\
M'_{24} &= \begin{bmatrix} m_{26} \\ 0 \end{bmatrix}, \\
M'_{32} &= \begin{bmatrix} m_{42} & m_{43} \\ m_{52} & m_{53} \end{bmatrix}, \\
M'_{33} &= \begin{bmatrix} m_{44} & 0 \\ 0 & m_{55} \end{bmatrix}, \\
M'_{42} &= \begin{bmatrix} m_{62} \\ 0 \end{bmatrix}, \\
M_{23} &= \begin{bmatrix} M'_{23} & M'_{24} \end{bmatrix}, \\
M_{32} &= \begin{bmatrix} M'_{32} \\ M'_{42} \end{bmatrix}, \\
M_{33} &= \begin{bmatrix} M'_{33} & 0 \\ 0 & 0 \end{bmatrix} \quad (A23)
\end{aligned}$$

in which the following conditions should hold in order to guarantee the negative value of M

- (i) $-2Q < 0$,
- (ii) $-2I_2 < 0$,
- (iii) Shur complement for $-I_2$ is $D_{-I_2} = -I_2 - M_{23} M_{33}^{-1} M_{32} < 0$ with proper parameter κ_i , $i = 1, 2, 3, 4$,
- (iv) Shur complement for M_{33} is $D_{M_{33}} = M_{33} + M_{32} I_2^{-1} M_{23} < 0$ with proper matrix κ_i , $i = 1, 2, 3, 4$.

Equation (A20) satisfies the following relationship

$$\dot{L} < -\|\tilde{L}\|^2 \sigma_{\min}(M) + \tilde{M}_{\max} \|\tilde{L}\| + \bar{M}_{\max} + \bar{\varepsilon}_1 + \bar{\varepsilon}_2 \quad (A24)$$

and derives negative Lyapunov derivative by completing the squares

$$\|\tilde{L}\| > -\frac{\tilde{M}_{\max}}{2\sigma_{\min}(M)} + \sqrt{\frac{\tilde{M}_{\max}^2}{4\sigma_{\min}^2(M)} - \frac{\bar{M}_{\max} + \bar{\varepsilon}_1 + \bar{\varepsilon}_2}{\sigma_{\min}(M)}} \quad (A25)$$

which implies that $\dot{L}$ is negative if the norm $\|\tilde{L}\|$ exceeds a certain bound.

The above result is obtained using the standard Lyapunov extension theorem, which demonstrates that both the state and the weights are uniformly ultimately bounded.

The error dynamics for $\hat{W}_1$ and $\hat{W}_2$ are given by

$$\begin{aligned}\dot{\tilde{W}}_1 &= -\alpha_1 \bar{\sigma}_{\hat{u}}^T \bar{\sigma}_{\hat{u}} \tilde{W}_1 + \alpha_1 \bar{\sigma}_{\hat{u}} \frac{\varepsilon_{HJ_1}}{m_{s1}} + \frac{\alpha_1}{m_{s1}^2} y^T A^T \nabla \phi_2^T(x) \times \\ &\quad \tilde{W}_2 + \frac{\alpha_1}{4m_{s1}^2} \left[\tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u + 2W_1^T \bar{E}_2(x) \tilde{W}_w \right] \quad (A26)\end{aligned}$$

$$\begin{aligned}\dot{\tilde{W}}_2 &= -\alpha_2 \bar{\sigma}_{\hat{w}}^T \bar{\sigma}_{\hat{w}} \tilde{W}_2 + \alpha_2 \bar{\sigma}_{\hat{w}} \frac{\varepsilon_{HJ_2}}{m_{s2}} + \\ &\quad \frac{\alpha_2}{4m_{s2}^2} \left[\tilde{W}_w^T \bar{D}_2(x) \tilde{W}_w + \tilde{W}_u^T \bar{D}_1(x) \tilde{W}_u + \right. \\ &\quad \left. 2\tilde{W}_u^T \bar{E}_1(x) W_2 - 2\tilde{W}_u^T \bar{D}_1(x) W_1 \right] \quad (A27)\end{aligned}$$

The following gives the proof of effective bounds of $\|\bar{\sigma}_{\hat{u}} \tilde{W}_1\|$ and $\|\bar{\sigma}_{\hat{w}} \tilde{W}_2\|$

$$\begin{aligned}\|\bar{\sigma}_{\hat{u}} \tilde{W}_1\| &> \varepsilon_{\max_1} > \bar{\varepsilon}_1 \left\| \frac{1}{m_{s1}} \right\| + \|y\| \left\| \frac{A^T \nabla \phi_2^T(x)}{m_{s1}} \right\| \|\tilde{W}_2\| + \\ &\quad \frac{1}{4} \|\tilde{W}_u\|^2 \left\| \frac{\bar{D}_1(x)}{m_{s1}} \right\| + \frac{1}{2} \|W_1\| \times \\ &\quad \left\| \frac{\bar{E}_2(x)}{m_{s1}} \right\| \|\tilde{W}_w\| \quad (A28)\end{aligned}$$

$$\begin{aligned}\|\bar{\sigma}_{\hat{w}} \tilde{W}_2\| &> \varepsilon_{\max_2} > \bar{\varepsilon}_2 \left\| \frac{1}{m_{s2}} \right\| + \frac{1}{4} \|\tilde{W}_w\|^2 \left\| \frac{\bar{D}_2(x)}{m_{s2}} \right\| + \frac{1}{4} \times \\ &\quad \|\tilde{W}_u\|^2 \left\| \frac{\bar{D}_1(x)}{m_{s2}} \right\| + \frac{1}{2} \|W_2\| \left\| \frac{\bar{E}_1(x)}{m_{s2}} \right\| \times \\ &\quad \|\tilde{W}_u\| - \frac{1}{2} \|W_1\| \left\| \frac{\bar{D}_1(x)}{m_{s2}} \right\| \|\tilde{W}_u\| \quad (A29)\end{aligned}$$

Thus, the proof is complete. $\blacksquare$

APPENDIX B

PROOF OF THEOREM 2

Proof The approximate coupled HJ equations are

$$\begin{aligned}H_2(x, \hat{W}_2, \hat{w}', \hat{u}') &= \\ H_2(x, \hat{W}_1, \hat{W}_2) &= \\ x^T Q x + \frac{1}{4} \hat{W}_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) \hat{W}_1 + \\ &\quad \frac{1}{4\gamma^2} \hat{W}_2^T \nabla \phi_2(x) B_2 B_2^T \nabla \phi_2^T(x) \hat{W}_2 + \hat{W}_2^T \times \\ &\quad \nabla \phi_2(x) A x - \frac{1}{2} \hat{W}_2^T \nabla \phi_2(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) \times \\ &\quad \hat{W}_1 - \varepsilon_{HJ_2} \quad (A30)\end{aligned}$$

$$\begin{aligned}H_1(x, \hat{W}_1, \hat{w}', \hat{u}') &= \\ H_1(x, \hat{W}_1, \hat{W}_2) &= \\ x^T Q x - \frac{1}{4} \hat{W}_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) \hat{W}_1 + \\ &\quad \hat{W}_1^T \nabla \phi_1(x) A x + \frac{1}{2\gamma^2} \hat{W}_1^T \nabla \phi_1(x) B_2 B_2^T \nabla \phi_2^T(x) \hat{W}_2 - \\ &\quad y^T (A^T \nabla \phi_2^T(x) \hat{W}_2 + 2Qx) - \varepsilon_{HJ_1} \quad (A31)\end{aligned}$$

Combining Eqs. (36) and (37), Eqs. (A30) and (A31) can be rewritten as

$$\begin{aligned}H_2(x, \hat{W}_1, \hat{W}_2) &= \\ \tilde{W}_2^T \nabla \phi_2(x) A x - \frac{1}{4} \hat{W}_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) \hat{W}_1 + \\ &\quad \frac{1}{4} W_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x) W_1 + \frac{1}{2} \hat{W}_2^T \nabla \phi_2(x) \times \\ &\quad B_1 R^{-1} B_1^T \nabla \phi_1^T(x) \hat{W}_1 - \frac{1}{2} W_2^T \nabla \phi_2(x) B_1 R^{-1} \times \\ &\quad B_1^T \nabla \phi_1^T(x) W_1 - \frac{1}{4\gamma^2} \tilde{W}_2^T \nabla \phi_2(x) B_2 B_2^T \nabla \phi_2^T(x) \tilde{W}_2 + \\ &\quad \frac{1}{2\gamma^2} W_2^T \nabla \phi_2(x) B_2 B_2^T \nabla \phi_2^T(x) \tilde{W}_2 - \varepsilon_{HJ_2} \quad (A32)\end{aligned}$$

$$\begin{aligned}H_1(x, \hat{W}_1, \hat{W}_2) &= \\ \tilde{W}_1^T \nabla \phi_1(x) A x - \frac{1}{2\gamma^2} \hat{W}_1^T \nabla \phi_1(x) B_2 B_2^T \nabla \phi_2^T(x) \hat{W}_2 + \\ &\quad \frac{1}{2\gamma^2} W_1^T \nabla \phi_1(x) B_2 B_2^T \nabla \phi_2^T(x) W_2 - \frac{1}{2} W_1^T \nabla \phi_1(x) \times \\ &\quad B_1 R^{-1} B_1^T \nabla \phi_1^T(x) \tilde{W}_1 + \frac{1}{4} \tilde{W}_1^T \nabla \phi_1(x) B_1 R^{-1} B_1^T \times \\ &\quad \nabla \phi_1^T(x) \tilde{W}_1 - y^T A^T \nabla \phi_2^T(x) \tilde{W}_2 - \varepsilon_{HJ_1} \quad (A33)\end{aligned}$$

Letting $\|W_1\| \leq W_{1,\max}$ and $\|W_2\| \leq W_{2,\max}$, thus we have

$$\begin{aligned}\|\hat{W}_1\| &= \|\tilde{W}_1 + W_1\| \leq \\ &\quad \|\tilde{W}_1\| + \|W_1\| \leq \|\tilde{W}_1\| + W_{1,\max} \quad (A34)\end{aligned}$$

$$\begin{aligned}\|\hat{W}_2\| &= \|\tilde{W}_2 + W_2\| \leq \\ &\quad \|\tilde{W}_2\| + \|W_2\| \leq \|\tilde{W}_2\| + W_{2,\max} \quad (A35)\end{aligned}$$

By taking into account Eqs. (A34) and (A35) and $\sup_{x \in \Omega} \|\varepsilon_{HJ_1}\| \leq \bar{\varepsilon}_1$ and $\sup_{x \in \Omega} \|\varepsilon_{HJ_2}\| \leq \bar{\varepsilon}_2$, Eqs. (A32) and (A33) become

$$\begin{aligned}\|H_2(x, \hat{W}_1, \hat{W}_2)\| &\leq \\ b_{\phi_2, x} \|A x\| \|\tilde{W}_2\| + \frac{1}{4} (\|\tilde{W}_1 + W_{1,\max}\|)^2 \|\nabla \phi_1(x) B_1\| \times \\ &\quad R^{-1} B_1^T \nabla \phi_1^T(x) + \frac{1}{4} W_{1,\max}^2 \|\nabla \phi_1(x) B_1 R^{-1} B_1^T \times \\ &\quad \nabla \phi_1^T(x)\| + \frac{1}{2} (\|\tilde{W}_1 + W_{1,\max}\|)(\|\tilde{W}_2 + W_{2,\max}\|) \times \\ &\quad \|\nabla \phi_2(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x)\| + \frac{1}{2} W_{1,\max} W_{2,\max} \times \\ &\quad \|\nabla \phi_2(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x)\| + \frac{1}{4\gamma^2} \|\tilde{W}_2\|^2 \|\nabla \phi_2(x) \times \\ &\quad B_2 B_2^T \nabla \phi_2^T(x)\| + \frac{1}{2\gamma^2} W_{2,\max} \|\tilde{W}_2\| \times \\ &\quad \|\nabla \phi_2(x) B_2 B_2^T \nabla \phi_2^T(x)\| + \bar{\varepsilon}_2 \quad (A36)\end{aligned}$$

$$\begin{aligned}\|H_1(x, \hat{W}_1, \hat{W}_2)\| &= \\ b_{\phi_1, x} \|A x\| \|\tilde{W}_1\| + \frac{1}{2\gamma^2} (\|\tilde{W}_1 + W_{1,\max}\|)(\|\tilde{W}_2 + \\ &\quad W_{2,\max}\|) \|\nabla \phi_1(x) B_2 B_2^T \nabla \phi_2^T(x)\| + \frac{1}{2\gamma^2} W_{1,\max} \times \\ &\quad W_{2,\max} \|\nabla \phi_1(x) B_2 B_2^T \nabla \phi_2^T(x)\| - \frac{1}{2} W_{1,\max} \|\tilde{W}_1\| \times \\ &\quad \|\nabla \phi_1(x) B_1 R^{-1} B_1^T \nabla \phi_1^T(x)\| + \frac{1}{4} \|\tilde{W}_1\|^2 \|\nabla \phi_1(x) \times \\ &\quad B_1 R^{-1} B_1^T \nabla \phi_1^T(x)\| - \|\tilde{W}_2\| b_{\phi_2, x} \|A y\| + \bar{\varepsilon}_1 \quad (A37)\end{aligned}$$

The terms on the right-hand side of Eqs. (A36) and (A37) are uniformly ultimately bounded, so both $H_2(x, \hat{W}_2, \hat{w}'(x), \hat{u}'(x))$ and $H_1(x, \hat{W}_1, \hat{w}'(x), \hat{u}'(x))$ are uniformly ultimately bounded. Furthermore, $\hat{W}_1$ and $\hat{W}_2$ converge to an approximate coupled Hamilton-Jacobi solution.

Based on Theorems 1 and 2 and Eqs. (59)–(62), it follows that $\|\hat{w} - w'\|$ and $\|\hat{u} - u'\|$ are uniformly ultimately bounded, since $\|\hat{W}_w - W_2\|$ and $\|\hat{W}_u - W_1\|$ are uniformly ultimately bounded.

Therefore, the policy pair $(\hat{w}(x), \hat{u}(x))$ provides an approximate equilibrium solution for the Stackelberg game.

This completes the proof. ■

ACKNOWLEDGMENT

This work was supported by the Support Plan on Science and Technology for Youth Innovation of Universities in Shandong Province (No. 2021KJ086), the National Natural Science Foundation of China (Nos. 62003234 and 61873179), and the National Natural Science Foundation of Shandong Province (No. ZR2020QF048).

REFERENCES

- [1] T. Başar and G. J. Olsder, *Dynamic Noncooperative Game Theory*, 2nd Ed. Philadelphia, PA, USA: SIAM, 1999.
- [2] M. O. Rabin, Effective computability of winning strategies, in *Contributions to the Theory of Games*. Princeton, NJ, USA: Princeton University Press, 1957, 147–157.
- [3] U. Schwalbe and P. Walker, Zermelo and the early history of game theory, *Games Econ. Behav.*, 2001, 34(1), 123–137.
- [4] S. M. LaValle and S. A. Hutchinson, Optimal motion planning for multiple robots having independent goals, *IEEE Trans. Rob. Autom.*, 1998, 14(6), 912–925.
- [5] Z. J. Ma, D. S. Callaway, and I. A. Hiskens, Decentralized charging control of large populations of plug-in electric vehicles, *IEEE Trans. Control Syst. Technol.*, 2013, 21(1), 67–78.
- [6] D. J. N. Limebeer, B. D. O. Anderson, and B. Hendel, A Nash game approach to mixed H_2/H_∞ control, *IEEE Trans. Autom. Control*, 1994, 39(1), 69–82.
- [7] K. G. Vamvoudakis, F. L. Lewis, and W. E. Dixon, Open-loop Stackelberg learning solution for hierarchical control problems, *Int. J. Adapt. Control Signal Process.*, 2019, 33(2), 285–299.
- [8] M. Li, J. H. Qin, Q. C. Ma, W. X. Zheng, and Y. Kang, Hierarchical optimal synchronization for linear systems via reinforcement learning: A Stackelberg-Nash game perspective, *IEEE Trans. Neural Netw. Learn. Syst.*, 2021, 32(4), 1600–1611.
- [9] I. Atzeni, L. G. Ordóñez, G. Scutari, D. P. Palomar, and J. R. Fonollosa, Noncooperative and cooperative optimization of distributed energy generation and storage in the demand-side of the smart grid, *IEEE Trans. Signal Process.*, 2013, 61(10), 2454–2472.
- [10] C. Dextreit and I. V. Kolmanovsky, Game theory controller for hybrid electric vehicles, *IEEE Trans. Control Syst. Technol.*, 2014, 22(2), 652–663.
- [11] X. Q. Li, W. Wang, J. J. Xu, and H. S. Zhang, LQG control on mixed H_2/H_∞ problem: The discrete-time case, *Int. J. Syst. Sci.*, 2020, 51(1), 191–201.
- [12] X. Q. Li, J. J. Xu, W. Wang, and H. S. Zhang, Mixed H_2/H_∞ control for discrete-time systems with input delay, *IET Control Theory Appl.*, 2018, 12(16), 2221–2231.
- [13] J. J. Xu, H. S. Zhang, and T. Y. Chai, Necessary and sufficient condition for two-player Stackelberg strategy, *IEEE Trans. Autom. Control*, 2015, 60(5), 1356–1361.
- [14] R. Z. Song, F. L. Lewis, and Q. L. Wei, Off-policy integral reinforcement learning method to solve nonlinear continuous-time multiplayer nonzero-sum games, *IEEE Trans. Neural Netw. Learn. Syst.*, 2017, 28(3), 704–713.
- [15] A. Odekinle, W. N. Gao, M. Davari, and Z. P. Jiang, Reinforcement learning and non-zero-sum game output regulation for multi-player linear uncertain systems, *Automatica*, 2020, 112, 108672.
- [16] Q. C. Zhang and D. B. Zhao, Data-based reinforcement learning for nonzero-sum games with unknown drift dynamics, *IEEE Trans. Cybern.*, 2019, 49(8), 2874–2885.
- [17] H. Kebriai and L. Iannelli, Discrete-time robust hierarchical linear-quadratic dynamic games, *IEEE Trans. Autom. Control*, 2018, 63(3), 902–909.
- [18] D. Vrabie and F. Lewis, Neural network approach to continuous-time direct adaptive optimal control for partially unknown nonlinear systems, *Neural Netw.*, 2009, 22(3), 237–246.
- [19] F. L. Lewis, D. Vrabie, and K. G. Vamvoudakis, Reinforcement learning and feedback control: Using natural decision methods to design optimal adaptive controllers, *IEEE Control Syst. Mag.*, 2012, 32(6), 76–105.
- [20] F. L. Lewis and D. Vrabie, Reinforcement learning and adaptive dynamic programming for feedback control, *IEEE Circ. Syst. Mag.*, 2009, 9(3), 32–50.
- [21] H. Modares, F. Lewis, and M. B. Naghibi-Sistani, Adaptive optimal control of unknown constrained-input systems using policy iteration and neural networks, *IEEE Trans. Neural Netw. Learn. Syst.*, 2013, 24(10), 1513–1525.
- [22] H. Modares and F. L. Lewis, Optimal tracking control of nonlinear partially-unknown constrained-input systems using integral reinforcement learning, *Automatica*, 2014, 50(7), 1780–1792.
- [23] P. Ioannou and B. Fidan, *Advances in Design and Control: Adaptive Control Tutorial*. Philadelphia, PA, USA: SIAM, 2006.
- [24] K. G. Vamvoudakis and F. L. Lewis, Online actor-critic algorithm to solve the continuous-time infinite horizon optimal control problem, *Automatica*, 2010, 46(5), 878–888.



Xiaoqian Li received the ME degree in control engineering from Northeastern University, Shenyang, China, in 2013, and the PhD degree in control science and engineering from Shandong University, Jinan, China, in 2019. She is currently an associate professor at Taishan University, Taian, China. Her research interests include optimal control, reinforcement learning, and game theory.



Zonglei Jing received the ME degree from Beihang University, Beijing, China, in 2024. He is currently pursuing the PhD degree at State Key Laboratory of Software Development Environment, Beihang University, China. His research interests include reinforcement learning, adversarial robustness, and optimal control.



Peijun Ju received the BE degree in mathematics from Qufu Normal University, Jining, China, in 2000, and the ME degree in mathematics and the PhD degree in control science and engineering from Shandong University, Jinan, China, in 2004 and 2017, respectively. He is currently a professor at Taishan University, China. His research interests include time delay systems, optimal control, and game theory.