

New Solution for H-infinity Static Output-Feedback Control Using Integral Reinforcement Learning

Yusuf Kartal^a, Wenqian Xue^b, Ahmet Taha Koru^a, Frank L. Lewis^a, Atilla Dogan^a

^a*University of Texas at Arlington, UTA Research Institute, Fort Worth, 76118, TX, USA*

^b*Northeastern University, State Key Laboratory of Synthetical Automation for Process Industries and International Joint Research Laboratory of Integrated Automation, Shenyang, 110819, China*

Abstract

This paper presents a new formulation for the H_∞ static output-feedback (OPFB) control problem that guarantees stability, L_2 gain boundedness, and optimal stabilizing gain solutions for a Linear Time-Invariant (LTI) system. Then, based on the developed formulation, it reveals an integral reinforcement learning algorithm (IRL) that employs Kleinman's method and achieves stabilizing optimal control strategies without requiring any knowledge of the system state, control, and disturbance matrices. The standard approaches to the H_∞ static OPFB control problem result in sub-optimal and non-Nash gain solutions for guaranteed stability. They propose an offline algorithm that may converge to only a sub-optimal stabilizing gain solution. On the other hand, this paper proposes a novel augmented Hamiltonian functional to solve the global Nash equilibrium solutions for a game of this kind. Based on the augmented Hamiltonian's stationarity conditions, we provide novel necessary and sufficient conditions for Nash equilibrium gain solutions that inherently stabilize the system dynamics while also guaranteeing L_2 gain bound by a prescribed attenuation level. To obtain the Nash solution without knowledge of system parameters, two off-policy IRL algorithms are developed based on Kleinman's algorithm. In the first IRL algorithm, the convergence to the Nash gain solution is provided assuming system state data is available. Then, a second novel IRL algorithm is developed that does not require system state data and learns the Nash equilibrium gain solution online. Simulation results are provided to show the validity of the proposed methods.

Keywords: H_∞ optimal control, integral reinforcement learning, zero-sum game, bounded L_2 gain

1. Introduction

One of the primary objectives in control system design is often to seek a stabilizing controller to regulate the output of a system that experiences disturbances. However, stability is not the only requirement in control system design. An L_2 gain bound of the system, optimality of the control method, and detectability of unknown system parameters are other common design specifications. Existing solutions of H_∞ static output-feedback (OPFB) yield stability and bounded L_2 gain, but employee non-Nash equilibrium solutions. In the two-player zero-sum game context, the Nash equilibrium consists of optimal strategies for both players. Based on this, the new formulation of H_∞ static OPFB control method is developed in this paper, which is a key to meet these requirements since it guarantees L_2 gain bound of the system by a prescribed attenuation level, asymptotic stability of equilibrium point, and also Nash equilibrium solutions. Then, based on the new formulae, the IRL algorithm is developed that iteratively solves H_∞ static OPFB Lyapunov equation, and deals with unknown system parameters.

H_∞ control methods has been widely studied in the literature, [1], [2], [3], [4], [5], [6] due to their applicability in variety of engineering areas. Some of these methods guarantee stability and bounded L_2 gain, but an extra condition is required to yield Nash equilibrium. [1] uses this method to design a gain-scheduled normal acceleration control loop for an air-launched unmanned aerial vehicle. Authors of [2] apply this control method on an industrial-type mass spring damper system. The efficacy of control law and the disturbance accommodation properties are shown on a rotor-craft design example in [7]. Moreover, [8] develop an autopilot controller for an F-16 aircraft by using the H_∞ static OPFB control method on a linear discrete-time system.

During the last few years, reinforcement learning (RL) algorithms [9], [10] has been used extensively to replace model parameters with a collected system's data [11], [12], [13], [14], [15], [16]. Particularly, offline iterative RL algorithms were studied in [17], [18], and [19]. The work by [17] considers the two-player policy iterations to solve for the feedback strategies of a continuous-time zero-sum game [20] in a sub-optimal manner that requires complete knowledge of system parameters. [12] presented an online RL

algorithm to solve the linear quadratic tracking (LQT) problem for partially-unknown continuous-time systems. In [21], the authors prove the convergence of IRL algorithm to a sub-optimal OPFB solution without considering the disturbance term when the drift dynamics are unknown. The optimal average cost learning framework is introduced to solve the output regulation problem for linear systems with unknown dynamics is studied in [22].

To design an efficient RL algorithm and achieve data-driven optimal control, the RL-based controller designs with neural networks (NNs) in an actor-critic structure [23], critic-only form [24] are proposed. In the off-policy RL algorithm, the system data, which are used to learn the solution of the corresponding Hamiltonian, can be generated with arbitrary policies rather than the evaluating policy. This approach is suitably implemented using NNs in [25] and [13].

The standard existing solutions to the static OPFB regulator problem [1], [21], [26], [27], [28] require some additive gain matrix to prove the stability of equilibrium, origin. Unfortunately, this results in the non-Nash equilibrium solutions. Consequently, the IRL algorithms developed based on these approaches [28], result in sub-optimal, non-Nash solutions. In this paper, we propose a novel augmented Hamiltonian, and develop new iterative algorithms based on stationarity conditions of the augmented Hamiltonian to obtain Nash equilibrium solutions. The salient contributions of this paper are summed up into four categories:

- This paper presents a new solution of the H_∞ static OPFB control problem. This solution guarantees not only stability and bounded L_2 gain but also Nash equilibrium solutions considering the corresponding min-max game.
- The existing methods employ an offline algorithm that does not have a convergence guarantee. Further, if the algorithm converges, it only converges sub-optimal stabilizing gain solutions. Our new solution rigorously yields Nash gain solutions given the novel necessary and sufficient conditions.
- To develop an IRL algorithm based on the augmented Hamiltonian's stationarity conditions, two off-line iterative solution algorithms are given. The first algorithm is based on Kleinman's algorithm and updates the disturbance gain term. A second algorithm modifies Kleinman's algorithm and gets the first algorithm in the IRL applicable form

that only updates the control gain.

- Two off-policy IRL algorithms are developed based on the modified Kleinman's algorithm. The first IRL algorithm learns the Nash gain solution online without requiring any knowledge of system dynamics' state, control, and disturbance matrices assuming the system state data is available. The second IRL algorithm assumes only the output data is available and develops a model-free observer to learn the Nash gain solution.

The rest of the paper is organized as follows. In Section 2, preliminaries on control design requirements are introduced, and the formulation of H_∞ static OPFB control is presented. A new solution of optimal H_∞ control problem and corresponding offline iterative solution algorithms are given in Section 3. In Section 4, an online off-policy IRL algorithm is developed based on stationarity conditions obtained in Section 3. Finally, we have shown the effectiveness of the proposed algorithms by applying them to the linearized lateral dynamics of the F-16 aircraft at a particular flight condition in Section 5.

Notations. We use the following notations throughout this paper $\mathbf{I}_n \in \mathbb{R}^{n \times n}$ is the identity matrix. The condition $A > 0$ (≥ 0) denotes the positive (semi) definiteness of a matrix. The operator $tr()$ denotes trace of a matrix. $\mathbf{C}^+ = \mathbf{C}^T(\mathbf{C}\mathbf{C}^T)^{-1}$ is the right-inverse of the full row-rank matrix $\mathbf{C}$ and the Kronecker product operator is denoted by $\otimes$. The determinant of a square matrix is denoted by $|\cdot|$. $vec(A)$ stands for the mn -vector formed by stacking the columns of $A \in \mathbb{R}^{n \times m}$ on top of one another, i.e., $vec(A) = [a_1^T \dots a_m^T]^T$ where $a_i \in \mathbb{R}^n$ are the columns of A . Lastly, $diag(\zeta_i)$ represents a diagonal matrix with $\zeta_i \forall i \in 1, \dots, N$ on its diagonal.

2. Preliminaries and problem formulation

In this section, preliminaries on Linear Time Invariant (LTI) system, and the corresponding controller design requirements are first introduced. Then, the problem description is presented.

2.1. System description and definitions

This section introduces system dynamics and performance specifications that are of interest. Consider the state-space representation of the continuous-

time LTI system as

$$\begin{aligned}\dot{\mathbf{x}} &= \mathbf{A}\mathbf{x} + \mathbf{B}\mathbf{u} + \mathbf{D}\mathbf{d} \\ \mathbf{y} &= \mathbf{C}\mathbf{x}\end{aligned}\tag{1}$$

where $\mathbf{A} \in \mathbb{R}^{n \times n}$, $\mathbf{B} \in \mathbb{R}^{n \times m}$, $\mathbf{D} \in \mathbb{R}^{n \times p}$ are system-state, input, disturbance matrices, and $\mathbf{C} \in \mathbb{R}^{q \times n}$ is assumed to be a full row-rank output matrix to avoid redundant measurements. The corresponding vectors $\mathbf{x}(t)$, $\mathbf{u}(t)$, $\mathbf{d}(t)$, and $\mathbf{y}(t)$ stand for the state, input, disturbance and output respectively.

Assumption 1. The pair $(\mathbf{A}, \mathbf{B})$ is *stabilizable* and the pair $(\mathbf{A}, \mathbf{C})$ is *detectable*.

Assumption 2. The system (1) is OPFB stabilizable because the row-space of output matrix $\mathbf{C}$ contains the sub-space spanned by the right eigenvectors corresponding to the unstable modes of $\mathbf{A}$.

Assumption 3. The non-zero columns of the output matrix $\mathbf{C}$ are linearly independent.

Remark 1. The Assumption 2 can be interpreted such that all unstable modes are measured by the output matrix $\mathbf{C}$ that represents the sensors installed in the system (1). The Assumption 3 enables us to recover a state element x_i precisely from the output vector $\mathbf{y}$ once it is left multiplied with $\mathbf{C}^+$.

Now, define the fictitious performance output $\mathbf{z}(t)$ that satisfies

$$\|\mathbf{z}\|_2^2 = \mathbf{x}^T \mathbf{Q} \mathbf{x} + \mathbf{u}^T \mathbf{R} \mathbf{u}\tag{2}$$

where $\mathbf{Q} \geq 0$ and $\mathbf{R} > 0$ are symmetric design matrices with appropriate dimensions. We assume that $\mathbf{Q}$ is selected such that the pair $(\mathbf{A}, \sqrt{\mathbf{Q}})$ is *observable*, which is a standard assumption [29]. Using the property $\|\mathbf{d}\|_2^2 = \mathbf{d}^T \mathbf{d}$, a realization of the following inequality $\forall \mathbf{d} \in [0, \infty)$ implies that the system L_2 -gain is bounded by a prescribed disturbance attenuation level denoted by γ

$$\int_0^\infty \|\mathbf{z}\|_2^2 dt \leq \gamma^2 \int_0^\infty \|\mathbf{d}\|_2^2 dt + \beta\tag{3}$$

for any non-zero energy-bounded disturbance input $\mathbf{d}$ [30] where β is a non-negative constant. The condition (3) is also called as non-expansivity constraint in [31]. Call γ^* the minimum gain for which this occurs. In [32] and

[33], an algorithm to find γ^* is given, and a formulation for explicit γ^* that depends on Riccati equation solution is derived for LTI systems under some assumptions. This paper assumes that the attenuation level is prescribed and satisfies $\gamma > \gamma^*$.

A static OPFB control to regulate the system (1) is

$$\mathbf{u} = -\mathbf{K}\mathbf{y} = -\mathbf{K}\mathbf{C}\mathbf{x} \quad (4)$$

where $\mathbf{K} \in \mathbb{R}^{m \times q}$ is the gain matrix. Note that main objective of H_∞ control using OPFB is to find the stabilizing $\mathbf{K}$ in an optimal manner while satisfying the condition (3), which can be achieved by solving corresponding Hamilton-Jacobi-Isaacs (HJI) equation.

2.2. Problem formulation and existing solution of static OPFB Problem

In this section, we relate zero-sum differential game theory to the static OPFB regulation problem in a global optimal manner by revealing various definitions. To satisfy L_2 -gain bound (3) with the stabilizing gain in (4), an objective functional defined as

$$J(\mathbf{u}, \mathbf{d}) = \int_0^\infty (\mathbf{x}^T \mathbf{Q} \mathbf{x} + \mathbf{u}^T \mathbf{R} \mathbf{u} - \gamma^2 \mathbf{d}^T \mathbf{d}) d\tau. \quad (5)$$

Now H_∞ control problem can be represented as a two-player zero-sum differential game by treating $\mathbf{u}(t)$ as a minimizing player, whereas $\mathbf{d}(t)$ maximizing player of (5). Then, the game can be formulated as

$$V(\mathbf{x}(0)) = J(\mathbf{u}^*, \mathbf{d}^*) = \min_{\mathbf{u}} \max_{\mathbf{d}} J(\mathbf{u}, \mathbf{d}) \quad (6)$$

where $V(\mathbf{x})$ denotes the value functional corresponding to (5) such that

$$V(\mathbf{x}) = \int_t^\infty (\mathbf{x}^T \mathbf{Q} \mathbf{x} + \mathbf{u}^T \mathbf{R} \mathbf{u} - \gamma^2 \mathbf{d}^T \mathbf{d}) d\tau. \quad (7)$$

and the pair $(\mathbf{u}^*, \mathbf{d}^*)$ denotes the game theoretic saddle point. The game of this kind admits a unique solution pair $(\mathbf{u}^*, \mathbf{d}^*)$, if the following Nash condition holds

$$\min_{\mathbf{u}} \max_{\mathbf{d}} J(\mathbf{u}, \mathbf{d}) = \max_{\mathbf{d}} \min_{\mathbf{u}} J(\mathbf{u}, \mathbf{d}). \quad (8)$$

The next theorem recalls the necessary and sufficient conditions for the sub-optimal H_∞ OPFB control method [1].

Theorem 1. *The system (1) is OPFB stable using the control $\mathbf{u}_o^e = -\mathbf{K}_o^e \mathbf{y}$ with L_2 gain bounded by $\gamma > \gamma^*$ if and only if*

1. *$(\mathbf{A}, \mathbf{B})$ is stabilizable and $(\mathbf{A}, \mathbf{C})$ is detectable.*
2. *There exists $\mathbf{K}_o^e$ and $\mathbf{L}$ such that*

$$\mathbf{K}_o^e = \mathbf{R}^{-1}(\mathbf{B}^T \mathbf{P} + \mathbf{L})\mathbf{C}^+ \quad (9)$$

where $\mathbf{P} = \mathbf{P}^T > \mathbf{0}$ is the solution of Riccati equation

$$\mathbf{P}\mathbf{A} + \mathbf{A}^T \mathbf{P} - \mathbf{P}\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \mathbf{P} + \mathbf{Q} + \gamma^{-2}\mathbf{P}\mathbf{D}\mathbf{D}^T \mathbf{P} + \mathbf{L}^T \mathbf{R}^{-1} \mathbf{L} = \mathbf{0}. \quad (10)$$

Proof. See [1] and [30] for the same proof. $\square$

Remark 2. On the one hand, introducing the additive gain matrix $\mathbf{L}$ provides the extra design freedom. Note that if $\mathbf{L} = \mathbf{0}$ there may not be a stabilizing solution to (9),(10) (See Theorem 1 in [1]). On the other hand, this results in a sub-optimal solution for the gain $\mathbf{K}^s$ (9). The Nash equilibrium gain occurs if the Theorem 1 holds with $\mathbf{L} = \mathbf{0}$.

The following lemma recalls the Nash equilibrium solution for the standard H_∞ control problem.

Lemma 1. The pair $(\mathbf{u}^*, \mathbf{d}^*)$ constitutes the Nash equilibrium of the game (8) such that

$$\mathbf{u}^* = -\mathbf{K}^* \mathbf{x} = -\mathbf{R}^{-1} \mathbf{B}^T \mathbf{P} \mathbf{x}, \quad (11)$$

$$\mathbf{d}^* = \gamma^{-2} \mathbf{D}^T \mathbf{P} \mathbf{x}. \quad (12)$$

Proof. Begin with deriving the Hamiltonian to solve the game theoretic saddle point or Nash equilibrium strategy of the game (8) as

$$H(\mathbf{V}_x, \mathbf{u}, \mathbf{d}) = \mathbf{x}^T \mathbf{Q} \mathbf{x} + \mathbf{u}^T \mathbf{R} \mathbf{u} - \gamma^2 \mathbf{d}^T \mathbf{d} + \mathbf{V}_x (\mathbf{A} \mathbf{x} + \mathbf{B} \mathbf{u} + \mathbf{D} \mathbf{d}) \quad (13)$$

where $\mathbf{V}_x = \partial V / \partial \mathbf{x}$ is the co-state vector. Using the quadratic form $V(\mathbf{x}) = \mathbf{x}^T \mathbf{P} \mathbf{x}$, and applying the stationarity conditions $\partial H(\mathbf{V}_x, \mathbf{u}, \mathbf{d}) / \partial \mathbf{u} = \mathbf{0}$ and $\partial H(\mathbf{V}_x, \mathbf{u}, \mathbf{d}) / \partial \mathbf{d} = \mathbf{0}$ yields the optimal control and disturbance respectively as (11) and (12).

Notice that the sign of Hessians, $\partial^2 H(\mathbf{V}_x, \mathbf{u}, \mathbf{d})/\partial \mathbf{u}^2 > 0$ and $\partial^2 H(\mathbf{V}_x, \mathbf{u}, \mathbf{d})/\partial \mathbf{d}^2 < 0$, along with unboundedness of the limits $\lim_{d \rightarrow \infty} J(\mathbf{u}^*, \mathbf{d})$, $\lim_{u \rightarrow \infty} J(\mathbf{u}, \mathbf{d}^*)$ indeed show that (11) and (12) are the global optimal minimizing and maximizing extrema respectively [34]. This indeed verifies that the pair $(\mathbf{u}^*, \mathbf{d}^*)$ denotes the Nash equilibrium point, which completes the proof. $\square$

Remark 3. The HJI equation, $H(\mathbf{V}_x, \mathbf{u}^*, \mathbf{d}^*) = 0$ with the boundary condition $V(0) = 0$ can be obtained by substituting the expressions (11)-(12) into Hamiltonian (13), which also verifies the sub-optimality of gain expression (9). Additionally, from the HJI equation $H(\mathbf{V}_x, \mathbf{u}^*, \mathbf{d}^*) = 0$, the Game Algebraic Riccati Equation (GARE) can be obtained as

$$PA + A^T P - PBR^{-1}B^T P + Q + \gamma^{-2}PDD^T P = 0. \quad (14)$$

If there is no $\mathbf{K}^s$ to satisfy (9) and (10), the OPFB H_∞ control problem may not have even a sub-optimal solution. In the next section, we rigorously analyze this and reveal some novel results.

3. New solution of H_∞ OPFB Game

Notice that Theorem 1 provides necessary and sufficient conditions for static OPFB in a sub-optimal manner. Thence, this does not yield a Nash equilibrium solution unless $L = 0$. Moreover, there may not even exist a static OPFB solution to the equations in Theorem 1. In this main section, two methods are proposed to solve H_∞ OPFB problem. The first method parameterizes the state feedback gains by using the Nash strategies and applies them to the OPFB design. The second method derives the new optimal H_∞ OPFB regulator formulation by introducing an augmented Hamiltonian. This method is introduced by [29], but is highly overlooked in the literature. However, it appears to be instrumental in H_∞ regulator design.

3.1. Necessary and Sufficient Conditions for the Stabilizing Nash Gain

Herein, we first parameterize static state feedback gains and then explain how to apply them to the H_∞ OPFB design.

Theorem 2. *Given the necessary conditions in the Assumption 1 and the sufficient condition $BR^{-1}B^T \geq \gamma^{-2}DD^T$. The system (1) is asymptotically stable using the control $\mathbf{u}^* = -R^{-1}B^T P \mathbf{x}$ (11) with $d = 0$, and L_2 gain bounded by $\gamma \forall \|\mathbf{d}\|_2 \in (0, \infty)$.*

Proof. To prove L_2 gain bound condition (3), first re-write the Hamiltonian (13) by completing the squares as

$$H(\mathbf{V}_x, \mathbf{u}, \mathbf{d}) = H(\mathbf{V}_x, \mathbf{u}^*, \mathbf{d}^*) + (\mathbf{u} - \mathbf{u}^*)^T \mathbf{R}(\mathbf{u} - \mathbf{u}^*) - \gamma^{-2}(\mathbf{d} - \mathbf{d}^*)^T(\mathbf{d} - \mathbf{d}^*). \quad (15)$$

Then, the objective functional can be re-expressed as

$$J(\mathbf{u}, \mathbf{d}) = \int_0^\infty \left(H(\mathbf{V}_x, \mathbf{u}^*, \mathbf{d}^*) + (\mathbf{u} - \mathbf{u}^*)^T \mathbf{R}(\mathbf{u} - \mathbf{u}^*) - \gamma^{-2}(\mathbf{d} - \mathbf{d}^*)^T(\mathbf{d} - \mathbf{d}^*) \right) dt + V(\mathbf{x}(0)). \quad (16)$$

Realize that the non-expansivity constraint (3) implies that $J(\mathbf{u}, \mathbf{d}) \leq \beta$. Select $\beta = V(\mathbf{x}(0))$, $\mathbf{u} = \mathbf{u}^*$, and note that HJI equation $H(\mathbf{V}_x, \mathbf{u}^*, \mathbf{d}^*) = 0$ holds with the boundary condition $V(0)=0$. Then, (16) reduces to

$$J(\mathbf{u}^*, \mathbf{d}) = - \int_0^\infty \gamma^{-2}(\mathbf{d} - \mathbf{d}^*)^T(\mathbf{d} - \mathbf{d}^*) dt + V(\mathbf{x}(0)), \leq V(\mathbf{x}(0)), \quad \forall \|\mathbf{d}\|_2 \in (0, \infty). \quad (17)$$

This proves that the L_2 gain bound condition (3) holds with $\beta = V(\mathbf{x}(0))$, $\mathbf{u} = \mathbf{u}^*$. Additionally, the value of game (8) with $\mathbf{u} = \mathbf{u}^*$ and $\mathbf{d} = \mathbf{d}^*$ is $V(\mathbf{x}(0))$ by (16).

To verify asymptotic stability, consider Lyapunov function $V(\mathbf{x}) = \mathbf{x}^T \mathbf{P} \mathbf{x}$ where $\mathbf{P}$ is solution of the GARE (14). Note that for $\mathbf{B} \mathbf{R}^{-1} \mathbf{B}^T \geq \gamma^{-2} \mathbf{D} \mathbf{D}^T$, the GARE (14) has a unique stabilizing solution $\mathbf{P} = \mathbf{P}^T$, which is indeed positive definite. This is illustrated in the following realization

$$\begin{aligned} \dot{V} &= \dot{\mathbf{x}}^T \mathbf{P} \mathbf{x} + \mathbf{x}^T \mathbf{P} \dot{\mathbf{x}} \\ &= \mathbf{x}^T (-\mathbf{P} \mathbf{B} \mathbf{R}^{-1} \mathbf{B}^T \mathbf{P} - \mathbf{Q} + \gamma^{-2} \mathbf{P} \mathbf{D} \mathbf{D}^T \mathbf{P}) \mathbf{x} \\ &\leq -\mathbf{x}^T \mathbf{Q} \mathbf{x} \quad \Leftarrow \quad \mathbf{B} \mathbf{R}^{-1} \mathbf{B}^T \geq \gamma^{-2} \mathbf{D} \mathbf{D}^T. \end{aligned} \quad (18)$$

Then, the observability of (A, Q) verifies that the undisturbed system (1), i.e., $d = 0$, is asymptotically stable by LaSalle's invariance principle [31]. This completes the proof. $\square$

Corollary 1. To verify asymptotic stability of the disturbed system, benefit from gain margin $[c_{lower}, \infty)$ with $c_{lower} < \frac{1}{2}$ property of the $\mathcal{H}_\infty$ control [31].

Note that the $\mathcal{H}_\infty$ has gain margin less than $\frac{1}{2}$ by Chapter 10 in [31] but the lower bound c_{lower} is not precisely defined. Then, if the sufficient condition in Theorem 2 is strengthened as $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2\gamma^{-2}\mathbf{D}\mathbf{D}^T$, the disturbed system (1) with $\mathbf{d} = \mathbf{d}^*$, becomes asymptotically stable. Thence, the closed-loop matrix $\mathbf{A} - \mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T\mathbf{P} + \gamma^{-2}\mathbf{D}\mathbf{D}^T\mathbf{P}$ becomes Hurwitz.

Till now, we actually parameterize all the stabilizing static state-feedback gains since we used the Nash strategies (11) and (12) in the proof of Theorem 2. Note that for the OPFB design, the Assumptions 2 and 3 are missing in the papers [1], [2], [7], and [30]. Therefore, the offline algorithm given in these works do not have a guaranteed convergence. In this manuscript, we use Assumptions 2 and 3 to apply state feedback gains to the OPFB design. The following Remark explains how to apply static state feedback gains to the H_∞ OPFB design.

Remark 4. Instead of Nash strategies (11) and (12), assume that the control and disturbance are selected as

$$\mathbf{u}_o^* = -(\mathbf{R}^{-1}\mathbf{B}^T\mathbf{P}\mathbf{C}^+)\mathbf{C}(\mathbf{C}^+\mathbf{y}) = -\underbrace{(\mathbf{R}^{-1}\mathbf{B}^T\mathbf{P}\mathbf{C}^+)}_{\mathbf{K}_o^*}\mathbf{y} \quad (19)$$

$$\mathbf{d}_o^* = (\gamma^{-2}\mathbf{D}^T\mathbf{P}\mathbf{C}^+)\mathbf{C}(\mathbf{C}^+\mathbf{y}). \quad (20)$$

Note that if $\mathbf{C}$ is an invertible matrix, then all states would be regulated optimally by Theorem 2. On the other hand, if it is not square but full row-rank, then only the states spanned by row space of the output matrix $\mathbf{C}$ would be regulated optimally given the Assumption 3, and the fact that $\mathbf{C}^+\mathbf{C}$ projects $\mathbb{R}^n$ onto the row space of $\mathbf{C}$. Additionally, the other states would converge to the origin given the Assumption 2. Realize that the Assumption 1 implies that there could be an unstable mode that is observable but does not belong to the row space of $\mathbf{C}$. Therefore, the Assumptions 2 and 3 are indeed required to apply static state feedback gains to the OPFB design. Lastly, the system (1) is stable against the worst-case disturbance (20), which affects only the states that belong to the row space of $\mathbf{C}$ if $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2\gamma^{-2}\mathbf{D}\mathbf{D}^T$ by Corollary 1.

3.2. A Direct Method to Obtain OPFB Optimal Gain Solutions

In this section, we propose a new methodology to obtain stabilizing Nash solutions for the H_∞ OPFB control. This method is direct in the sense that

it reaches the same gain solutions as the Section 3.1 but do not require two step

Consider the optimal value obtained in the Theorem 2 that corresponds to the zero-effort for each player such that

$$J_0 = \mathbf{x}^T(0)\mathbf{P}\mathbf{x}(0) = \text{tr}(\mathbf{P}\mathbf{X}_0) \quad (21)$$

where $\text{tr}()$ stands for the trace of a matrix, and $\mathbf{X}_0 = \mathbf{x}(0)\mathbf{x}^T(0)$.

The following Lemma is an essential step before introducing the augmented Hamiltonian.

Lemma 2. Given the Assumption 1, let $\mathbf{K}^a$ be a gain that stabilizes the system (1), and the corresponding OPFB control is $\mathbf{u}_o^a = -\mathbf{K}_o^a\mathbf{y}$. Additionally, let the disturbance takes the form $\mathbf{d}_o^a = -\mathbf{N}_o\mathbf{y}$ to guarantee it does not affect unobservable modes of the system (1). Then the corresponding Lyapunov equation can be derived as

$$\mathbf{P}\mathbf{A}_c + \mathbf{A}_c^T\mathbf{P} + \mathbf{C}^T\mathbf{K}_o^{aT}\mathbf{R}\mathbf{K}_o^a\mathbf{C} - \gamma^2\mathbf{C}^T\mathbf{N}_o^T\mathbf{N}_o\mathbf{C} + \mathbf{Q} \equiv \mathbf{T} \quad (22)$$

where $\mathbf{T} = \mathbf{T}^T = \mathbf{0}$ and $\mathbf{A}_c = \mathbf{A} - \mathbf{B}\mathbf{K}_o^a\mathbf{C} + \mathbf{D}\mathbf{N}_o\mathbf{C}$.

Proof. Consider the quadratic form of the value functional (5) as $V(\mathbf{x}) = \mathbf{x}^T\mathbf{P}\mathbf{x}$, and then substitute expressions $\mathbf{u} = -\mathbf{K}^a\mathbf{y}$ and $\mathbf{d}_o^a = -\mathbf{N}_o\mathbf{y}$ into (5) to obtain

$$\mathbf{x}^T\mathbf{P}\mathbf{x} = \int_t^\infty \mathbf{x}^T(\mathbf{C}^T\mathbf{K}_o^{aT}\mathbf{R}\mathbf{K}_o^a\mathbf{C} + \mathbf{Q} - \gamma^2\mathbf{C}^T\mathbf{N}_o^T\mathbf{N}_o\mathbf{C})\mathbf{x}d\tau. \quad (23)$$

Now, take the derivative of left-side (23) and substitute (12), $\mathbf{u} = -\mathbf{K}^a\mathbf{C}\mathbf{x}$ expressions. Lastly, take the derivative of integral in right-side (23) using Leibniz's rule, which yields

$$\begin{aligned} \mathbf{x}^T(\mathbf{A}_c^T\mathbf{P} + \mathbf{P}\mathbf{A}_c)\mathbf{x} = \\ \mathbf{x}^T(-\mathbf{C}^T\mathbf{K}_o^{aT}\mathbf{R}\mathbf{K}_o^a\mathbf{C} - \mathbf{Q} + \gamma^2\mathbf{C}^T\mathbf{N}_o^T\mathbf{N}_o\mathbf{C})\mathbf{x}. \end{aligned} \quad (24)$$

Realize that the zero equivalent is nothing but the Lyapunov equation given in (22). This completes the proof. $\square$

It is now clear that performing a min-max operation on (5), subject to dynamical constraint (1) is equivalent to the algebraic problem of finding the pair $(\mathbf{K}_o^a, \mathbf{N}_o)$ that performs min-max of (21) subject to the constraint (22).

Define the following augmented Hamiltonian to solve this modified problem as

$$H^a(\mathbf{K}_o^a, \mathbf{N}_o, \mathbf{S}) = \text{tr}(\mathbf{P}\mathbf{X}_0) + \text{tr}(\mathbf{T}\mathbf{S}) \quad (25)$$

where $\mathbf{S} \in \mathbb{R}^{n \times n}$ is a symmetric matrix of Lagrange multipliers [29] that needs to be determined, and $\mathbf{T}$ is given in (22). Notice that $\mathbf{T}$ includes weighting matrices $\mathbf{Q}$ and $\mathbf{R}$.

The next main theorem is a key to find $\mathbf{S}$ along with a Nash equilibrium control matrix $\mathbf{K}_o^a$ with respect to (25).

Theorem 3. *Given the Assumptions 1-3, the system (1) is asymptotically stable using the control $\mathbf{u}_o^a = -\mathbf{K}_o^a \mathbf{y}$ with $d_o^a = 0$ and L_2 gain bounded by γ if*

$$\frac{\partial H^a}{\partial \mathbf{S}} \equiv \mathbf{P}\mathbf{A}_c + \mathbf{A}_c^T \mathbf{P} + \mathbf{C}^T \mathbf{K}_o^{aT} \mathbf{R} \mathbf{K}_o^a \mathbf{C} - \gamma^2 \mathbf{C}^T \mathbf{N}_o^T \mathbf{N}_o \mathbf{C} + \mathbf{Q} = \mathbf{0}, \quad (26)$$

$$\frac{\partial H^a}{\partial \mathbf{P}} \equiv \mathbf{S}\mathbf{A}_c^T + \mathbf{A}_c \mathbf{S} + \mathbf{X}_0 = \mathbf{0}, \quad (27)$$

$$\frac{\partial H^a}{\partial \mathbf{K}_o^a} \equiv 2\mathbf{R}\mathbf{K}_o^a \mathbf{C} \mathbf{S} \mathbf{C}^T - 2\mathbf{B}^T \mathbf{P} \mathbf{S} \mathbf{C}^T = \mathbf{0}, \quad (28)$$

$$\frac{\partial H^a}{\partial \mathbf{N}_o} \equiv -2\gamma^2 \mathbf{N}_o^a \mathbf{C} \mathbf{S} \mathbf{C}^T + 2\mathbf{D}^T \mathbf{P} \mathbf{S} \mathbf{C}^T = \mathbf{0}. \quad (29)$$

Furthermore, the following gain expressions solves the H_∞ static OPFB problem in an optimal manner

$$\mathbf{K}_o^a = \mathbf{K}_o^* = \mathbf{R}^{-1} \mathbf{B}^T \mathbf{P} \mathbf{C}^+ \quad (30)$$

$$\mathbf{N}_o = \gamma^{-2} \mathbf{D}^T \mathbf{P} \mathbf{C}^+ \quad (31)$$

Proof. Consider (25), which is a constant along the system trajectories since the system (1) is LTI and $\mathbf{z}(t)$ in (2) is not explicit function of time. This implies that we can apply the constraint test and check the stationarity conditions on the augmented Hamiltonian (25) that yields the second-order Lyapunov equation (26) and standard Lyapunov equation (27) respectively.

Define a variable $\mathbf{X} = \mathbf{x}\mathbf{x}^T$ that includes the system state information. Taking the derivative of $\mathbf{X}$ using (1) yields

$$\begin{aligned} \dot{\mathbf{X}} &= \dot{\mathbf{x}}\mathbf{x}^T + \mathbf{x}\dot{\mathbf{x}}^T \\ &= \mathbf{A}_c \mathbf{x}\mathbf{x}^T + \mathbf{x}\mathbf{x}^T \mathbf{A}_c^T \\ &= \mathbf{A}_c \mathbf{X} + \mathbf{X} \mathbf{A}_c^T. \end{aligned} \quad (32)$$

Now, assume that $\mathbf{A}_c$ is Hurwitz. Then, taking the integral of both sides (32) from 0 to ∞ yields (27) where $S = \int_0^\infty X dt$, thereby $\mathbf{K}_o^a$ should not depend on the solution S . Therefore, the gain solutions (30) and (31) are immediate. Realize that the stability of an LTI system does not depend on the initial condition, i.e, local stability implies global stability. Thence, the gain solutions (30) and (31) should not depend on S that depends on the initial condition X_0 . This also verifies our reason to select gain solutions in the given forms, which completes the proof. $\square$

Remark 5. The Theorem 3 gives the necessary conditions, i.e. the Assumption 1 and 2, and sufficient condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq \gamma^{-2}\mathbf{D}\mathbf{D}^T$ to prove L_2 gain boundedness by a prescribed attenuation level γ (3) and OPFB stability considering the worst-case disturbance (20). Realize that the condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq \gamma^{-2}\mathbf{D}\mathbf{D}^T$ is only a sufficient condition, which implies that there may be an optimal gain solution which stabilizes (1) but does not satisfy $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq \gamma^{-2}\mathbf{D}\mathbf{D}^T$. However, in that case, one may not achieve a positive definite solution P for the Riccati equation (10). Additionally, the positive definiteness of the Riccati equation solution plays a key role for Kleinman's algorithm in Section 3.3, and the IRL algorithm in Section 4.

Remark 6. The Theorem 3 proves that the condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq \gamma^{-2}\mathbf{D}\mathbf{D}^T$ is *sufficient* to obtain the stabilizing Nash equilibrium solution. The gain $\mathbf{K}_o^e$ in Theorem 1 is a sub-optimal stabilizing gain solution with respect to the value functional (7). However, the gain solution K^a in Theorem 3 always gives a stabilizing Nash equilibrium gain solutions with respect to the game (8) given the Assumptions 1-3.

Remark 7. Note that the system (1) is stable in the presence of matched disturbances with the gains (30) and (31). The unmatched disturbances are only L_2 gain bounded by Theorem 2. Thence, the control (19) is robustly stabilizing the equilibrium origin even if the unmatched disturbances exist given Assumptions 2 and 3.

The next section proposes an offline model-based algorithm to find the optimal gain solution $\mathbf{K}_o^a$ iteratively, that plays a key role to develop the IRL Algorithm that will be detailed later in Section 4. Note that given the necessary and sufficient conditions in Theorem 3 and Remark 5, one does not need an iterative solution algorithm to find optimal stabilizing gain $\mathbf{K}_o^a$ (30). However, to develop a model-free algorithm, an iterative solution algorithm is required.

3.3. Offline Iterative Solution Algorithms for H_∞ OPFB

This section presents two iterative solution algorithms to obtain minimizing gain (30) by using the conditions given in (26)-(28). In the first algorithm, we employ Kleinman's algorithm (26) to obtain Nash equilibrium gain solutions (30) and (31), whereas in the second algorithm, we use a corresponding Lyapunov equation to not deal with the disturbance gain term $\mathbf{N}_o$.

The next algorithm performs a sequence of four-step iterations based on the Kleinman's Algorithm [35] to find the optimal control gains (30) and (31).

Algorithm 1. (Offline iterative solution with Lyapunov equations. Kleinman's Algorithm.)

1. Initialize: Set $k = 1$, $\mathbf{P}_0 = \mathbf{0}$, $\mathbf{N}_0 = \mathbf{0}$ and given the Assumption 1 and the sufficient condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2\gamma^{-2}\mathbf{D}\mathbf{D}^T$, select a gain $\mathbf{F}_0$ such that $\mathbf{A} - \mathbf{B}\mathbf{F}_0$ is asymptotically stable.
2. k^{th} iteration: Solve for $\mathbf{P}_k$

$$\mathbf{0} = \mathbf{P}_k\mathbf{A}_k + \mathbf{A}_k^T\mathbf{P}_k + \mathbf{F}_{k-1}^T\mathbf{R}\mathbf{F}_{k-1} + \mathbf{Q} \quad (33)$$

where $\mathbf{A}_k = \mathbf{A} - \mathbf{B}\mathbf{F}_k + \frac{1}{2}\mathbf{D}\mathbf{N}_k$. Finally, update the gains

$$\mathbf{F}_k = \mathbf{R}^{-1}\mathbf{B}^T\mathbf{P}_k, \quad (34)$$

$$\mathbf{N}_k = \gamma^{-2}\mathbf{D}^T\mathbf{P}_k. \quad (35)$$

Set the cost $J_k = \text{tr}(\mathbf{P}_k\mathbf{X}_0)$.

3. Check: If $\mathbf{F}_{k-1}$ and $\mathbf{F}_k$ are close enough to each other, go to step (4) else go to step (2).
4. Terminate: Given the Assumptions 2 and 3, set the OPFB gains $\mathbf{K}_o^a = \mathbf{F}_k\mathbf{C}^+$, $\mathbf{N}_o = \mathbf{N}_k\mathbf{C}^+$ and the cost $\mathbf{J} = \mathbf{J}_k$. $\square$

Note that the closed-loop stability, and L_2 gain boundedness implies that (36) has a unique stabilizing optimal solution $\mathbf{P} > \mathbf{0}$ by Theorem 2. A comprehensive study for the solution of generalized Riccati equations can be found in [36]. The Algorithm 2 is based on the iterative solution algorithm presented in [37] whose convergence is proved by establishing the connection

between Newton's method in [38]. Additionally, by considering the condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq \gamma^{-2}\mathbf{D}\mathbf{D}^T$, the monotonic convergence, i.e, $P_k < P_{k-1}$, is straight-forward from the Theorem 13.5.8 in [38]. A related algorithm is also examined in Chapter 8 of the book [29].

Now, to develop an algorithm that accounts only for the optimal gain K_o^a , we manipulate the steps of the Algorithm 1. The resultant Algorithm 2 finds the the optimal control gain K_o^a (30) iteratively.

Algorithm 2. (Offline iterative solution with Lyapunov equations. Modified Kleinman's Algorithm.)

1. Initialize: Set $k = 1$, $\mathbf{P}_0 = \mathbf{0}$, and given the Assumption 1 and the sufficient condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2\gamma^{-2}\mathbf{D}\mathbf{D}^T$, select a gain $\mathbf{F}_0$ such that $\mathbf{A} - \mathbf{B}\mathbf{F}_0$ is asymptotically stable.

2. k^{th} iteration: Solve for $\mathbf{P}_k$

$$\mathbf{0} = \mathbf{P}_k \mathbf{A}_k + \mathbf{A}_k^T \mathbf{P}_k + \mathbf{F}_{k-1}^T \mathbf{R} \mathbf{F}_{k-1} + \mathbf{Q} + \gamma^{-2} \mathbf{P}_{k-1} \mathbf{D} \mathbf{D}^T \mathbf{P}_{k-1} \quad (36)$$

where $\mathbf{A}_k = \mathbf{A} - \mathbf{B}\mathbf{F}_k$. Finally, update the gain

$$\mathbf{F}_k = \mathbf{R}^{-1} \mathbf{B}^T \mathbf{P}_k. \quad (37)$$

Set the cost $J_k = \text{tr}(\mathbf{P}_k \mathbf{X}_0)$.

3. Check: If $\mathbf{F}_{k-1}$ and $\mathbf{F}_k$ are close enough to each other, go to step (4) else go to step (2).
4. Terminate: Given the Assumptions 2 and 3, set the OPFB gain $\mathbf{K}_o^a = \mathbf{F}_k \mathbf{C}^+$ and the cost $\mathbf{J} = \mathbf{J}_k$. $\square$

The next section uses the Algorithm 2 to develop model-free algorithms considering different scenarios for the availability of system data.

4. Online Integral Reinforcement Learning Solution Algorithm for H_∞ OPFB

In this section, we first develop an online off-policy integral reinforcement learning (IRL) algorithm [39], which is a model-free version of the Algorithm 2. This algorithm assumes that the system state data is available to deal with unknown $\mathbf{A}$, $\mathbf{B}$, and $\mathbf{D}$ matrices. Then, we develop a novel IRL algorithm that solves the optimal H_∞ regulator problem by learning the Nash equilibrium gain solution (30) without requiring knowledge of the system state data.

4.1. Online off-policy IRL algorithms

This section introduces an off-policy IRL algorithm to deal with the unknown system matrices $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{D}$. In this case, both of Algorithm 1 and 2 lose their applicability as they are model-based. However, we still benefit from the convergence properties of Algorithm 2 while developing the off-policy IRL method. To this end, we represent system dynamics (1) as

$$\begin{aligned}\dot{\mathbf{x}} &= \mathbf{A}\mathbf{x} + \mathbf{B}\mathbf{u}_k + \mathbf{D}\mathbf{d} + \mathbf{B}(\mathbf{u} - \mathbf{u}_k) \\ &= \mathbf{A}_k\mathbf{x} + \mathbf{D}\mathbf{d} + \mathbf{B}(\mathbf{u} - \mathbf{u}_k)\end{aligned}\quad (38)$$

where $\mathbf{A}_k = \mathbf{A} - \mathbf{B}\mathbf{F}_k$ and $\mathbf{u}_k = -\mathbf{F}_k\mathbf{x} \in \mathbb{R}^m$ is the control policy to be updated with $\mathbf{F}_k$ given in Algorithm 2.

Firstly, to obtain $\mathbf{P}_k$ without information of the system matrices $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{D}$, take the derivative of value functional $V(\mathbf{x}(t)) = \mathbf{x}^T(t)\mathbf{P}\mathbf{x}(t)$ by using the new representation of the system dynamics (38)

$$\begin{aligned}\dot{V} &= \mathbf{x}^T \mathbf{A}_k^T \mathbf{P}_k \mathbf{x} + \mathbf{x}^T \mathbf{P}_k \mathbf{A}_k \mathbf{x} \\ &\quad + 2\mathbf{d}^T \mathbf{D}^T \mathbf{P}_k \mathbf{x} + 2(\mathbf{u} + \mathbf{F}_k \mathbf{x})^T \mathbf{B}^T \mathbf{P}_k \mathbf{x}.\end{aligned}\quad (39)$$

To employ the approach given in Algorithm [35], we define the following two new variables

$$\mathbf{G}_{k+1} = \gamma^{-2} \mathbf{D}^T \mathbf{P}_k, \quad \mathbf{F}_{k+1} = \mathbf{R}^{-1} \mathbf{B}^T \mathbf{P}_k. \quad (40)$$

Re-write (36) in terms of new variable (40) to get the Algorithm 2 in the Kleinman's form as

$$\bar{\mathbf{Q}} = \mathbf{P}_k \mathbf{A}_k + \mathbf{A}_k^T \mathbf{P}_k \quad (41)$$

where $\bar{\mathbf{Q}} = -\mathbf{F}_k^T \mathbf{R} \mathbf{F}_k - \mathbf{Q} - \gamma^2 \mathbf{G}_k^T \mathbf{G}_k$. Additionally, express (39) in terms of the new variables introduced in (40) and (41) as

$$\dot{V} = 2(\gamma^2 \mathbf{d}^T \mathbf{G}_{k+1} + (\mathbf{u} + \mathbf{F}_k \mathbf{x})^T \mathbf{R} \mathbf{F}_{k+1}) \mathbf{x} + \mathbf{x}^T \bar{\mathbf{Q}} \mathbf{x}. \quad (42)$$

Then, integrate both sides from t to $t + T$ to obtain

$$\begin{aligned}V(t + T) - V(t) &= \int_t^{t+T} 2(\mathbf{u} + \mathbf{F}_k \mathbf{x})^T \mathbf{R} \mathbf{F}_{k+1} \mathbf{x} d\tau \\ &\quad + \int_t^{t+T} 2\gamma^2 \mathbf{d}^T \mathbf{G}_{k+1} \mathbf{x} d\tau + \int_t^{t+T} \mathbf{x}^T \bar{\mathbf{Q}} \mathbf{x} d\tau.\end{aligned}\quad (43)$$

Based on these manipulations, the online off-policy IRL algorithm can be developed. Note that this new IRL algorithm and the Algorithm 2 are equivalent. However, on the contrary offline Algorithm 2, the IRL Algorithm 3 does not require information of $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{D}$ matrices. It only requires an initial stabilizing control policy $u_0 = -F_0x$ (this is standard assumption in IRL applications [14],[23],[25],[28]) and the sufficient amount of data that belongs $\mathbf{x}(t)$, $\mathbf{u}(t)$, and $\mathbf{d}(t)$ vectors, which can be collected online.

Algorithm 3. (Off-policy IRL algorithm assuming $\mathbf{x}$ is given.)

1. Initialize: Set $k = 1$, $\mathbf{G}_0 = \mathbf{0}$ and given Assumption 1, start with initial stabilizing control policy $u_0 = -F_0x$.
2. k^{th} iteration: Use (43) to update $\mathbf{P}_k$, $\mathbf{G}_{k+1}$ and $\mathbf{F}_{k+1}$ simultaneously

$$\begin{aligned} & \mathbf{x}^T(t+T)\mathbf{P}_k\mathbf{x}(t+T) - \mathbf{x}^T(t)\mathbf{P}_k\mathbf{x}(t) - \int_t^{t+T} 2\gamma^2 \mathbf{d}^T \mathbf{G}_{k+1} \mathbf{x} d\tau \\ & - \int_t^{t+T} 2(\mathbf{u} + \mathbf{F}_k \mathbf{x})^T \mathbf{R} \mathbf{F}_{k+1} \mathbf{x} d\tau = \int_t^{t+T} \mathbf{x}^T \bar{\mathbf{Q}} \mathbf{x} d\tau. \end{aligned} \quad (44)$$

Set the cost $J_k = \text{tr}(\mathbf{P}_k \mathbf{X}_0)$.

3. Check: If $\mathbf{F}_k$ and $\mathbf{F}_{k+1}$ are close enough to each other, go to step (4) else go to step (2).
4. Terminate: Given the Assumptions 2 and 3, set $\mathbf{K}_o^a = \mathbf{F}_{k+1} \mathbf{C}^+$ and $\mathbf{J} = \mathbf{J}_k$. $\square$

Remark 8. Note that right-side of the equation (44) in the Algorithm 3 consists of known terms, and hence it can be solved for $\mathbf{P}_k$, $\mathbf{G}_{k+1}$ and $\mathbf{F}_{k+1}$ matrices using well established least-squares technique by converting them to the set of linear equations [25]. Since (44) does not use any system matrix information except $\mathbf{C}$, the Algorithm 3 is said to be model-free. Realize that the output matrix $\mathbf{C}$ represents the sensors placed to the system (1), which is clearly known.

Realize that the Algorithm 3 achieves the static OPFB gain assuming system state information, $\mathbf{x}$, is available. On the other hand, if we are only given the output data $\mathbf{y}$, we need to come up with a novel algorithm that does not require system state data. One approach is to make use of the

observability matrix while developing a model-free algorithm. However, this approach indeed requires the pair $(\mathbf{A}, \mathbf{C})$ to be observable. Thence, from now on we assume that the detectability condition in the Assumption 1 is strengthened as an observability condition.

The next Theorem is instrumental before we develop a new Algorithm that only need measurement of the system output data $\mathbf{y}$.

Theorem 4. *For any given n -dimensional observable system, there exists a sufficiently small time interval $[0, T]$ such that if N sampling times satisfy $0 \leq t - i\Delta t \leq T \forall i \in 1, \dots, N$ where Δt is the delayed time and assumed fixed, then the system is N -sample observable.*

Proof.: See [40] for the same proof.

To make use of the observability matrix define

$$\mathbf{x}(t)^T \mathbf{P} \mathbf{x}(t) = \mathbf{Y}^T(t) \tilde{\mathbf{P}} \mathbf{Y}(t) \quad (45)$$

where $\mathbf{Y}(t) = [\mathbf{y}^T(t) \ \dot{\mathbf{y}}^T(t) \ \dots \ \mathbf{y}^{n-1T}(t)]^T = \mathbf{O} \mathbf{x}$, and hence $\tilde{\mathbf{P}} = \mathcal{O}_L^T \mathbf{P} \mathcal{O}_L$ with $\mathcal{O}_L \in \mathbb{R}^{n \times nq}$ is the left-inverse of the observability matrix $\mathbf{O} \in \mathbb{R}^{nq \times n}$. Note that each derivative of the output $\mathbf{y}$ can be obtained by making use of the Taylor Series expansion on the collected data $\mathbf{y}(t + i\Delta t)$ around $\mathbf{y}(t)$. An example is $\ddot{\mathbf{y}} = \frac{\mathbf{y}(t+\Delta t) + \mathbf{y}(t-\Delta t) - 2\mathbf{y}(t)}{\Delta t^2}$ where $\mathbf{y}(t + \Delta t) = \mathbf{y}(t) + \Delta t \dot{\mathbf{y}}(t) + 0.5\Delta^2 t \ddot{\mathbf{y}}(t)$ and $\mathbf{y}(t - \Delta t) = \mathbf{y}(t) - \Delta t \dot{\mathbf{y}}(t) + 0.5\Delta^2 t \ddot{\mathbf{y}}(t)$.

Now, select $\mathbf{Q} = k\mathbf{C}^T \mathbf{C}$ where $k > 0$ is a scalar, and define the following variables

$$\tilde{\mathbf{F}}_k = \mathbf{F}_k \mathcal{O}_L, \quad \tilde{\mathbf{G}}_k = \mathbf{G}_k \mathcal{O}_L, \quad (46)$$

and the known term $\tilde{\mathbf{Q}} = -\tilde{\mathbf{F}}_k^T \mathbf{R} \tilde{\mathbf{F}}_k - \hat{\mathbf{Q}} - \gamma^2 \tilde{\mathbf{G}}_k^T \tilde{\mathbf{G}}_k$. Herein the new weighting matrix selected in a form such that $\hat{\mathbf{Q}} = k[\mathbf{I}_q \ \mathbf{0} \ \dots \ \mathbf{0}]_{qn \times q}^T [\mathbf{I}_q \ \mathbf{0} \ \dots \ \mathbf{0}]_{q \times qn}$. Realize that $\mathbf{x}^T \mathbf{Q} \mathbf{x} = \mathbf{y}^T \mathbf{y} = \mathbf{Y}^T \hat{\mathbf{Q}} \mathbf{Y}$ holds when $\mathbf{Q} = k\mathbf{C}^T \mathbf{C}$, and it enables us to treat $\tilde{\mathbf{Q}}$ as a known term. Further, since the pair $(\mathbf{A}, \mathbf{C})$ is observable $(\mathbf{A}, k\mathbf{C}^T \mathbf{C})$ is also observable. Based on these manipulations, a new Algorithm 4 can be developed.

Algorithm 4. (Off-policy IRL algorithm, $\mathbf{x}$ is not required but the pair $(\mathbf{A}, \mathbf{C})$ must be observable.)

1. Initialize: Set $k = 1$, $\tilde{\mathbf{G}}_0 = \mathbf{0}$ and start with a stabilizing control policy $\mathbf{u} = -\tilde{\mathbf{F}}_0 \mathbf{Y}$.

2. k^{th} iteration: Update $\tilde{\mathbf{P}}_k$, $\tilde{\mathbf{G}}_{k+1}$ and $\tilde{\mathbf{F}}_{k+1}$ simultaneously

$$\begin{aligned} & \mathbf{Y}^T(t+T)\tilde{\mathbf{P}}_k\mathbf{Y}(t+T) - \mathbf{Y}^T(t)\tilde{\mathbf{P}}_k\mathbf{Y}(t) - \int_t^{t+T} 2\gamma^2 d^T \tilde{\mathbf{G}}_{k+1} \mathbf{Y} d\tau \\ & - \int_t^{t+T} 2(\mathbf{u} + \tilde{\mathbf{F}}_k \mathbf{Y})^T \mathbf{R} \tilde{\mathbf{F}}_{k+1} \mathbf{Y} d\tau = \int_t^{t+T} \mathbf{Y}^T \tilde{\mathbf{Q}} \mathbf{Y} d\tau. \end{aligned} \quad (47)$$

Set the cost $J_k = \text{tr}(\tilde{\mathbf{P}}_k \mathbf{X}_0)$.

3. Check: If $\tilde{\mathbf{F}}_k$ and $\tilde{\mathbf{F}}_{k+1}$ are close enough to each other, go to step (4) else go to step (2).
4. Terminate: Given the Assumptions 2 and 3, set $\mathbf{u} = -\tilde{\mathbf{F}}_{k+1} \mathbf{Y}$ and $\mathbf{J} = \mathbf{J}_k$. $\square$

Realize that the Algorithm 4 converges with the static state feedback expression since $\mathbf{u} = -\tilde{\mathbf{F}}_{k+1} \mathbf{Y} = -\mathbf{R}^{-1} \mathbf{B}^T \mathbf{P}_k \mathcal{O}_L \mathcal{O} \mathbf{x} = -\mathbf{R}^{-1} \mathbf{B}^T \mathbf{P}_k \mathbf{x}$ (11). Thence, it regulates not only the state elements spanned in the row-space of $\mathbf{C}$ but all states. However, it creates additional complexity as it requires more data to be collected to converge. On the other hand, the Algorithm 4 does not require system state data $\mathbf{x}$, and it directly achieves the Nash equilibrium control u by making use of the new variables $\tilde{\mathbf{F}}_{k+1}$ and $\mathbf{y}$ instead of calculating the OPFB gain $\mathbf{K}_o^a$. Additionally, since the Algorithm 4 is obtained with the change of variables in the Algorithm 3, it shares similar convergence properties with the Algorithm 3, and hence the Algorithm 2. The next section shows a way of solving coupled linear equations.

4.2. Data-based implementation of the IRL Algorithm

This section introduces a least-squares method to solve step (2) in Algorithm 3. Although value function approximation is a popular tool and can be employed to solve step 2 in Algorithm 3, it requires three Neural Networks (NNs), i.e., the actor NN to approximate the value functional (7), the critic NN to update control policy and the disturber NN to update disturbance policy [13]. This causes a complicated NN design procedure. Instead, we use a least-squares method to solve for $\mathbf{P}_k$ in step (2) of Algorithm 3, however, we first need to find $\mathbf{P}_k$.

Now, we use the Kronecker product property $\mathbf{a}^T \mathbf{W} \mathbf{b} = (\mathbf{b}^T \otimes \mathbf{a}^T) \text{vec}(\mathbf{W})$ to rewrite (44) as

$$\begin{aligned} & [\hat{\mathbf{x}}(t+T) - \hat{\mathbf{x}}(t)]^T \hat{\mathbf{P}}_k - 2\gamma^2 \left(\int_t^{t+T} \mathbf{x}^T \otimes \mathbf{d}^T d\tau \right) \text{vec}(\mathbf{G}_{k+1}) \\ & - 2 \left(\int_t^{t+T} \mathbf{x}^T \otimes [(\mathbf{u} + \mathbf{F}_k \mathbf{x})^T \mathbf{R}] d\tau \right) \text{vec}(\mathbf{F}_{k+1}) = \int_t^{t+T} \mathbf{x}^T \bar{\mathbf{Q}} \mathbf{x} d\tau \end{aligned} \quad (48)$$

where the vectors $\hat{\mathbf{x}} \in \mathbb{R}^{\frac{n(n-1)}{2}}$ and $\hat{\mathbf{P}}_k \in \mathbb{R}^{\frac{n(n-1)}{2}}$ are defined in the following forms

$$\begin{aligned} \hat{\mathbf{x}} &= [x_1^2, 2x_1x_2, \dots, 2x_1x_n, x_2^2, \dots, 2x_{n-1}x_n, x_n^2]^T \\ \hat{\mathbf{P}}_k &= [P_{k(11)}, \dots, P_{k(1n)}, P_{k(22)}, \dots, P_{k(2n)}, \dots, P_{k(nn)}]^T. \end{aligned} \quad (49)$$

To solve $\mathbf{P}_k$, $\mathbf{G}_{k+1}$ and $\mathbf{F}_{k+1}$ in (48), define

$$\mathbf{d}_x = [\hat{\mathbf{x}}(t+T) - \hat{\mathbf{x}}(t), \dots, \quad (50)$$

$$\hat{\mathbf{x}}(t+s_1T) - \hat{\mathbf{x}}(t+(s_1-1)T)]^T \in \mathbb{R}^{s_1 \times \frac{n(n-1)}{2}};$$

$$\begin{aligned} \mathbf{I}_{xd} &= \left[\int_t^{t+T} (\mathbf{x} \otimes \mathbf{d}) d\tau, \dots, \right. \\ & \left. \int_{t+(s_1-1)T}^{t+s_1T} (\mathbf{x} \otimes \mathbf{d}) d\tau \right]^T \in \mathbb{R}^{s_1 \times np}; \end{aligned} \quad (51)$$

$$\begin{aligned} \mathbf{I}_{xu} &= \left[\int_t^{t+T} \mathbf{x} \otimes (\mathbf{u} + \mathbf{F}_k \mathbf{x}) d\tau, \dots, \right. \\ & \left. \int_{t+(s_1-1)T}^{t+s_1T} \mathbf{x} \otimes (\mathbf{u} + \mathbf{F}_k \mathbf{x}) d\tau \right]^T \in \mathbb{R}^{s_1 \times nm}; \end{aligned} \quad (52)$$

$$\Phi = [\mathbf{d}_x, -2\mathbf{I}_{xd}, -2\mathbf{I}_{xu}(\mathbf{I}_n \otimes \mathbf{R})]; \quad (53)$$

$$\Psi = -\left[\int_t^{t+T} \mathbf{x}^T \bar{\mathbf{Q}} \mathbf{x} d\tau, \dots, \int_{t+(s_1-1)T}^{t+s_1T} \mathbf{x}^T \bar{\mathbf{Q}} \mathbf{x} d\tau \right]^T, \quad (54)$$

where integer $s_1 > 0$ is the sampling data group number. Then, the solution can be obtained by

$$\begin{bmatrix} \hat{\mathbf{P}}_k \\ \text{vec}(\mathbf{G}_{k+1}) \\ \text{vec}(\mathbf{F}_{k+1}) \end{bmatrix} = (\Phi^T \Phi)^{-1} \Phi^T \Psi. \quad (55)$$

Remark 9. To ensure that (55) achieves the unique solution, the persistence of the excitation condition needs to be satisfied. To this end, probing noise should be injected to control input u in (37). This is also called as exploration noise that does not affect the convergence [25]. In addition, the data group number s_1 should be no less than $\frac{n(n+1)}{2} + np + nm$, which is the number of unknown parameters to be calculated by (55).

In the next section, the correct performance of the proposed methods is validated by applying them to the lateral control of linearized F-16 dynamics.

5. Simulation Results

In this section, an example is given to verify the correct performance of proposed algorithms that solves optimal H_∞ static OPFB regulator problem. We used the F-16 linearized lateral dynamics at a particular flight condition given in example 5.3-1 of the book [41]. Parameters of the linearized F-16 lateral dynamics and the corresponding system vectors are

$$\begin{aligned} \mathbf{A} &= \begin{bmatrix} -0.32 & 0.064 & 0.036 & -0.992 & 0 & 0.001 \\ 0 & 0 & 1 & 0.004 & 0 & 0 \\ -30.65 & 0 & -3.68 & 0.665 & -0.733 & 0.132 \\ 8.54 & 0 & -0.025 & -0.476 & -0.032 & -0.062 \\ 0 & 0 & 0 & 0 & -20.2 & 0 \\ 0 & 0 & 0 & 0 & 0 & -20.2 \end{bmatrix}; \\ \mathbf{B} &= \begin{bmatrix} 0 & 0 & 0 & 0 & 20.2 & 0 \\ 0 & 0 & 0 & 0 & 0 & 20.2 \end{bmatrix}^T; \\ \mathbf{C} &= \begin{bmatrix} 0 & 0 & 0 & 57.2958 & 0 & 0 \\ 0 & 0 & 57.2958 & 0 & 0 & 0 \\ 57.2958 & 0 & 0 & 0 & 0 & 0 \\ 0 & 57.2958 & 0 & 0 & 0 & 0 \end{bmatrix}; \\ \mathbf{x} &= [\beta \ \phi \ p \ r \ \delta_a \ \delta_r]^T; \ \mathbf{u} = [u_a \ u_r]^T; \\ \mathbf{y} &= [r \ p \ \beta \ \phi]^T \end{aligned} \quad (56)$$

where β denotes the side-slip, ϕ is the bank angle, p is the roll rate, r is the yaw rate, δ_a is the aileron actuator angle, δ_r is the rudder actuator angle, u_a is the aileron servo input, u_r is the rudder servo input. The factor of 57.2958 in the output-matrix $\mathbf{C}$ converts radians to degrees.

In a real-life scenario, the system states, and disturbances can be measured through the sensors placed on the vehicle of interest. To exemplify, Inertial Measurement Units (IMUs) can be used to determine attitude (pitch-roll-yaw angles and their rates) of the vehicle, and structural health monitoring systems can be used to measure vibrational disturbances. In addition, the linearized model of F-16 (56) is only valid when the true speed is 502 *feet/sec* and 300 *psf* dynamic pressure. When the aircraft changes its flight condition, i.e., the true speed or dynamics pressure values change, the Algorithm 3 can be used to obtain stabilizing Nash control policy without knowing the new system parameters of the F-16 assuming the system states, $\mathbf{x}$, is available. If there is a sensor problem and all of the system states are not available, then the Algorithm 4 can be used since it only requires the system output data $\mathbf{y}$.

Additionally, the objective functional (5) parameters are selected as $\mathbf{Q} = \text{diag}([50 \ 100 \ 100 \ 50 \ 0 \ 0])$, $\mathbf{R} = \rho \times \text{diag}([0.1 \ 0.7])$ with $\rho = 0.3$ for computation of the OPFB gain. Also, select the disturbance matrix as

$$\mathbf{D} = \begin{bmatrix} 0 & 0 & 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0 & 0 & 2 \end{bmatrix}^T, \quad (57)$$

and set the attenuation level $\gamma = 2.5$. To examine the robustness, assume that the system (1) experiences the worst-case disturbance (20). Realize that all of the Assumptions 1-3, and the sufficient condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2\gamma^{-2}\mathbf{D}\mathbf{D}^T$ are satisfied with the given parameters in (56) and (57). Now, we first illustrate the performance of model-based Nash gain solutions (30) and (31), and then check whether the Algorithms 1 and 2 converges the same game solutions as (30) and (31). After verifying the correctness of them, we illustrate the performance of the Nash solutions obtained in the Algorithms 3 and 4 that are model-free.

The Fig. 1 illustrates the zero control-effort response of the system (1). The Fig. 2 illustrates the optimal stabilizing gain (30) performance, which is derived in Theorem 3. The Nash OPFB gains found by (30) and (31) are

$$\mathbf{K}_o^a = \begin{bmatrix} -0.1395 & -0.8714 & 1.4785 & -1.0000 \\ -0.1410 & 0.0273 & -0.1070 & 0.0330 \end{bmatrix}, \quad (58)$$

$$\mathbf{N}^o = \begin{bmatrix} -0.0001 & -0.0006 & 0.0011 & -0.0007 \\ -0.0005 & 0.0001 & -0.0004 & 0.0001 \end{bmatrix}. \quad (59)$$

Now, to examine the performance of the Algorithms 1 and 2, we set the

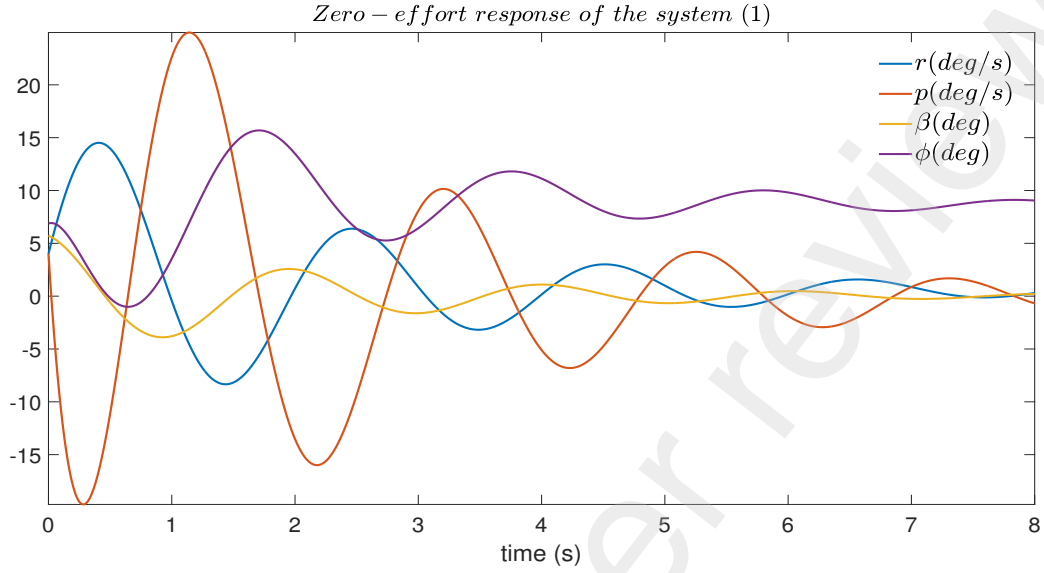


Figure 1: System response when no control is applied.

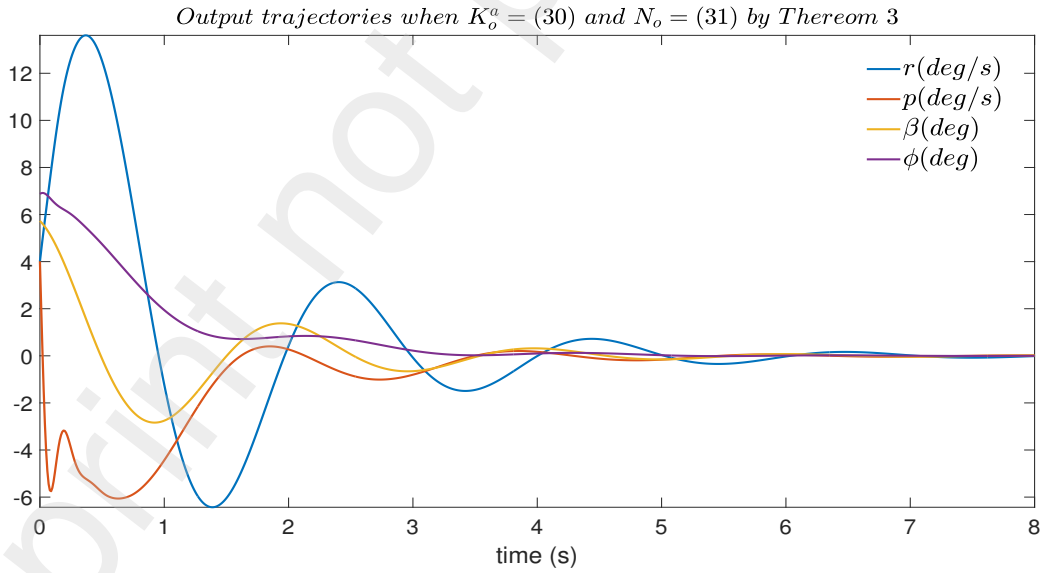


Figure 2: System response when the Nash gains (30) and (31) are employed.

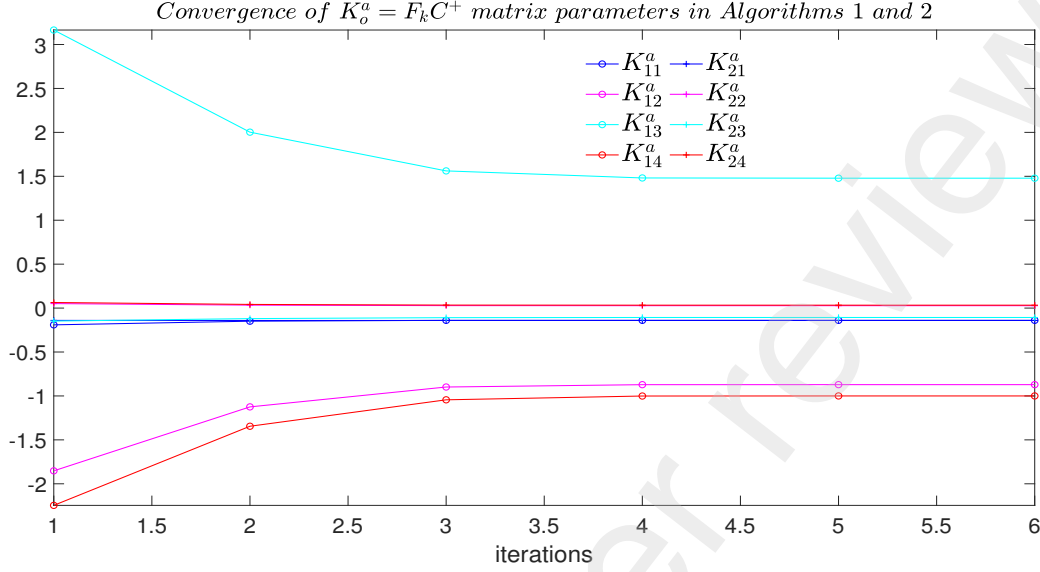


Figure 3: Convergence of the gain matrix $\mathbf{K}_o^a$ parameters by using the Algorithms 1 and 2.

initial sub-optimal stabilizing gain as

$$\mathbf{F}_0 = \begin{bmatrix} -0.0888 & -0.1875 & 0.7076 & -0.2328 \\ -0.1382 & 0.0105 & -0.0884 & 0.0141 \end{bmatrix} C. \quad (60)$$

The convergence of optimal gain matrix parameters $\mathbf{K}_o^a$ can be observed from Fig. 3. Note that their convergence properties are exactly the same as each other and they converge to the same gain matrix (58) as expected. Therefore, the output trajectories figure with this resultant optimal stabilizing gain K^a is the same as Fig. 2. Now compare the system responses in Fig. 1 and Fig. 2 to observe the regulation performance. The convergence of disturbance gain matrix parameters $\mathbf{N}_o$ is also illustrated in Fig. 4. Note that the converged parameters of $\mathbf{N}_o$ are the same as the OPFB disturbance gain given in (59) as illustrated.

Next, to show the correct performance of the model-free Algorithm 3, we set the initial stabilizing gain F_0 as the same as (60). Then, based on the state information gathered from the system, the Algorithm 3 is employed to learn the Nash equilibrium gain solution $\mathbf{K}_o^a$ online. Once the IRL Algorithm converged, we applied the corresponding control $\mathbf{u} = -\mathbf{K}_o^a \mathbf{y}$ using to the

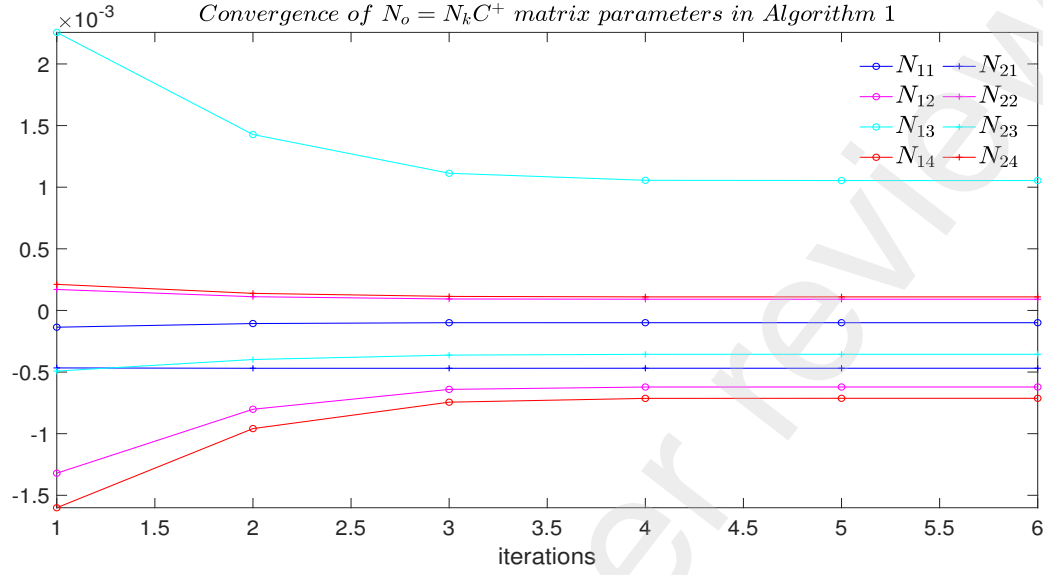


Figure 4: Convergence of the N_o matrix parameters by using the Algorithm 1.

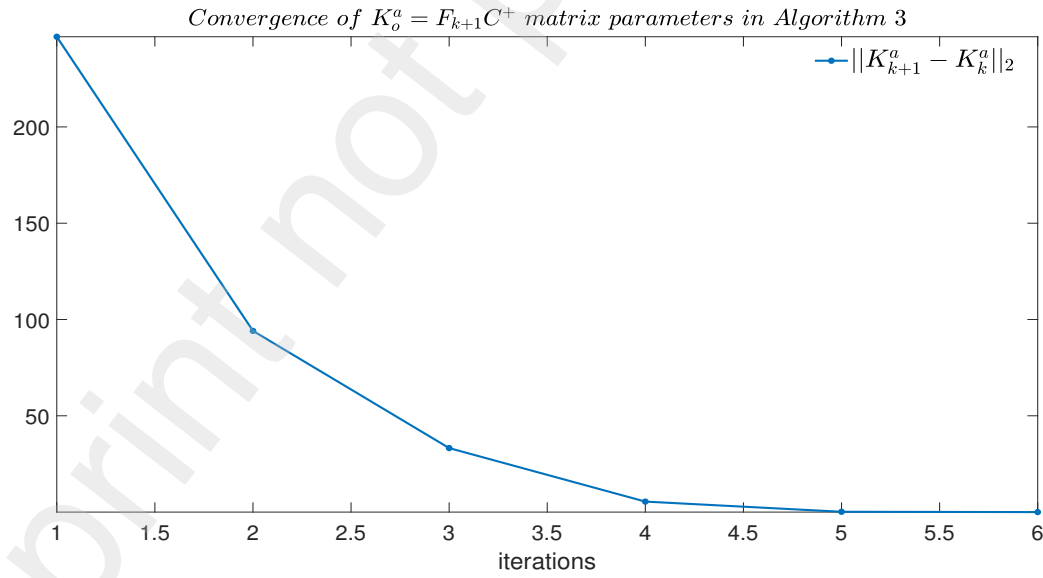


Figure 5: Convergence of the gain matrix K_o^a parameters by using the Algorithm 3.

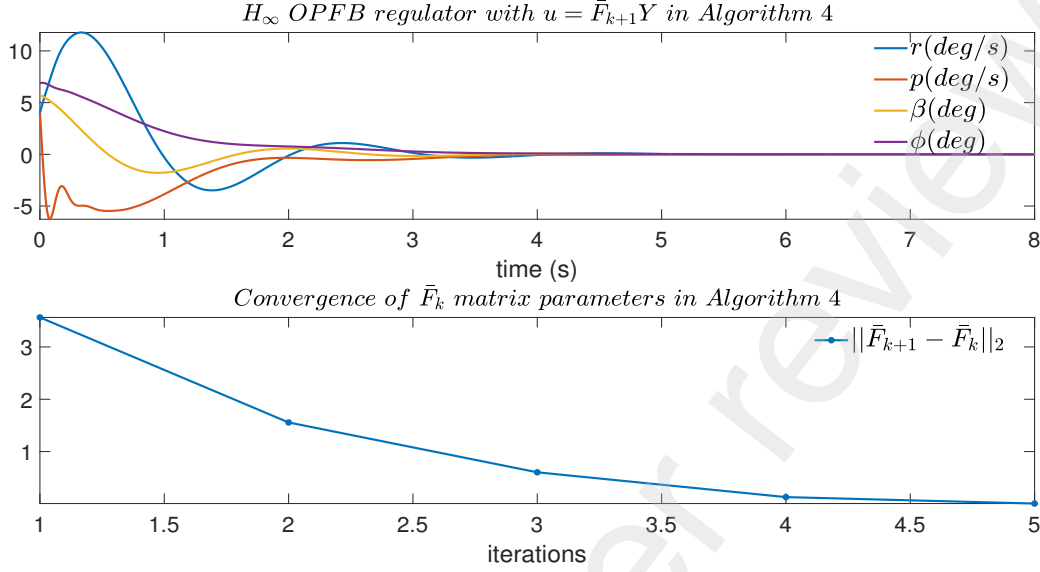


Figure 6: Performance of the Algorithm 4.

actual system (1). The Algorithm 3 is also converged to the same Nash gain $\mathbf{K}_o^a$ (58), thereby the resultant output vector states are the same as Fig. 2 as expected.

On the other hand, to check the correctness of the Algorithm 4, we select $\mathbf{Q} = k\mathbf{C}^T\mathbf{C}$ with $k = 0.05$ as explained in Section 4.1. The resultant output trajectories are shown in Fig. 6. Additionally, we have calculated the Nash state feedback gain expression $\mathbf{K}^*$ given in (11), and also the observability matrix $\mathcal{O}$ for the system (1). Then, we compared the $\mathbf{K}^*$ with the $\bar{\mathbf{F}}_k + 1\mathcal{O}$. It has been seen that the two matrices have almost the same elements and L_2 norm difference between them is calculated as 0.84, which verifies the correct performance of the Algorithm 4.

Lastly, the critical attenuation level obtained in the Theorems 2 and 3 is $\gamma^* = 0.05$. Note that with this critical attenuation γ^* level, the sufficient condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2(\gamma^*)^{-2}\mathbf{D}\mathbf{D}^T$ is relaxed. However, this condition indeed required for the Kleinman based Algorithms 1-4, and the critical attenuation level is obtained as $\gamma^* = 0.49$, which is also compatible with the sufficient condition $\mathbf{B}\mathbf{R}^{-1}\mathbf{B}^T \geq 2(\gamma^*)^{-2}\mathbf{D}\mathbf{D}^T$. Therefore, we conclude that the model-free algorithms reduce the L_2 gain performance.

6. Summary and Conclusions

In this paper, we proposed a novel augmented Hamiltonian, and develop new iterative algorithms based on stationarity conditions of the augmented Hamiltonian to obtain Nash equilibrium solutions. Additionally, the corresponding optimal gain solutions guarantee both stability and L_2 gain boundedness of an LTI system when the H_∞ static OPFB control method is employed. Convergence properties of two off-line iterative solution algorithms are given. Then, based on the Lyapunov iterations, an online off-policy IRL algorithm which is a model-free version of the offline Algorithm 2, is developed to solve the optimal H_∞ regulator problem by learning the Nash equilibrium solution (30) without requiring system state, control, and disturbance matrices. However, this Algorithm assumes the availability of the system state data. Thence, we come up with a novel model-free observer-based Algorithm 4 that only requires system output data to achieve stabilizing Nash control policies but it creates additional complexity as it requires more data to be collected to converge. Lastly, we applied the proposed algorithms to the linearized F-16 lateral dynamics at a particular flight condition to verify the correct performance of the proposed algorithms. Further research can be conducted to investigate how the state and output measurement delays affect the proposed methods.

7. Acknowledgements

This work is supported by the Office of Naval Research under Grant N00014-18-1-2221, the National Science Foundation under Grant ECCS-1839804 and Army Research Office under the Grant/Award Number W911NF-20-1-0132.

References

- [1] J. Gadewadikar, F. L. Lewis, M. Abu-Khalaf, Necessary and sufficient conditions for h -infinity static output-feedback control, *Journal of guidance, control, and dynamics* 29 (4) (2006) 915–920.
- [2] J. Gadewadikar, A. Bhilegaonkar, F. L. Lewis, Bounded l_2 gain static output feedback: Controller design and implementation on an electromechanical system, *IEEE Transactions On Industrial Electronics* 54 (5) (2007) 2593–2599.

- [3] A. Al-Tamimi, F. L. Lewis, M. Abu-Khalaf, Model-free q-learning designs for linear discrete-time zero-sum games with application to h-infinity control, *Automatica* 43 (3) (2007) 473–481.
- [4] S. A. A. Rizvi, Z. Lin, Output feedback q-learning for discrete-time linear zero-sum games with application to the h-infinity control, *Automatica* 95 (2018) 213–221.
- [5] Y. Jiang, K. Zhang, J. Wu, C. Zhang, W. Xue, T. Chai, F. L. Lewis, H-infinity based minimal energy adaptive control with preset convergence rate, *IEEE Transactions on Cybernetics* (2021).
- [6] J. Han, H. Zhang, H. Jiang, X. Sun, H-infinity consensus for linear heterogeneous multi-agent systems with state and output feedback control, *Neurocomputing* 275 (2018) 2635–2644.
- [7] J. Gadewadikar, F. L. Lewis, K. Subbarao, K. Peng, B. M. Chen, H-infinity static output-feedback control for rotorcraft, *Journal of Intelligent and Robotic Systems* 54 (4) (2009) 629–646.
- [8] A. P. Valadbeigi, A. K. Sedigh, F. L. Lewis, H infinity static output-feedback control design for discrete-time systems using reinforcement learning, *IEEE transactions on neural networks and learning systems* 31 (2) (2019) 396–406.
- [9] R. S. Sutton, A. G. Barto, *Reinforcement learning: An introduction*, MIT press, 2018.
- [10] D. P. Bertsekas, *Dynamic programming and optimal control: Vol. 1*, Athena scientific Belmont, 2000.
- [11] B. Kiumarsi, F. L. Lewis, H. Modares, A. Karimpour, M.-B. Naghibi-Sistani, Reinforcement q-learning for optimal tracking control of linear discrete-time systems with unknown dynamics, *Automatica* 50 (4) (2014) 1167–1175.
- [12] H. Modares, F. L. Lewis, Linear quadratic tracking control of partially-unknown continuous-time systems using reinforcement learning, *IEEE Transactions on Automatic control* 59 (11) (2014) 3051–3056.

- [13] H. Modares, F. L. Lewis, Z.-P. Jiang, H-infinity tracking control of completely unknown continuous-time systems via off-policy reinforcement learning, *IEEE transactions on neural networks and learning systems* 26 (10) (2015) 2550–2562.
- [14] B. Kiumarsi, K. G. Vamvoudakis, H. Modares, F. L. Lewis, Optimal and autonomous control using reinforcement learning: A survey, *IEEE transactions on neural networks and learning systems* 29 (6) (2017) 2042–2062.
- [15] C. Chen, H. Modares, K. Xie, F. L. Lewis, Y. Wan, S. Xie, Reinforcement learning-based adaptive optimal exponential tracking control of linear systems with unknown dynamics, *IEEE Transactions on Automatic Control* 64 (11) (2019) 4423–4438.
- [16] B. Luo, H.-N. Wu, T. Huang, Off-policy reinforcement learning for h-infinity control design, *IEEE transactions on cybernetics* 45 (1) (2014) 65–76.
- [17] M. Abu-Khalaf, F. L. Lewis, J. Huang, Neurodynamic programming and zero-sum games for constrained control systems, *IEEE Transactions on Neural Networks* 19 (7) (2008) 1243–1252.
- [18] H. Zhang, Q. Wei, D. Liu, An iterative adaptive dynamic programming method for solving a class of nonlinear zero-sum differential games, *Automatica* 47 (1) (2011) 207–214.
- [19] Q. Wei, D. Liu, Q. Lin, R. Song, Adaptive dynamic programming for discrete-time zero-sum games, *IEEE transactions on neural networks and learning systems* 29 (4) (2017) 957–969.
- [20] S. Mehraeen, T. Dierks, S. Jagannathan, M. L. Crow, Zero-sum two-player game theoretic formulation of affine nonlinear discrete-time systems using neural networks, *IEEE transactions on cybernetics* 43 (6) (2012) 1641–1655.
- [21] L. M. Zhu, H. Modares, G. O. Peen, F. L. Lewis, B. Yue, Adaptive sub-optimal output-feedback control for linear systems using integral reinforcement learning, *IEEE Transactions on Control Systems Technology* 23 (1) (2014) 264–273.

- [22] F. A. Yaghmaie, S. Gunnarsson, F. L. Lewis, Output regulation of unknown linear systems using average cost reinforcement learning, *Automatica* 110 (2019) 108549.
- [23] K. G. Vamvoudakis, F. L. Lewis, Online actor-critic algorithm to solve the continuous-time infinite horizon optimal control problem, *Automatica* 46 (5) (2010) 878–888.
- [24] H. Jiang, H. Zhang, X. Xie, Critic-only adaptive dynamic programming algorithms' applications to the secure control of cyber-physical systems, *ISA transactions* 104 (2020) 138–144.
- [25] Y. Jiang, Z.-P. Jiang, Computational adaptive optimal control for continuous-time linear systems with completely unknown dynamics, *Automatica* 48 (10) (2012) 2699–2704.
- [26] R. Moghadam, F. L. Lewis, Output-feedback h-infinity quadratic tracking control of linear systems using reinforcement learning, *International Journal of Adaptive Control and Signal Processing* 33 (2) (2019) 300–314.
- [27] Q. Jiao, H. Modares, F. L. Lewis, S. Xu, L. Xie, Distributed l2-gain output-feedback control of homogeneous and heterogeneous systems, *Automatica* 71 (2016) 361–368.
- [28] S. A. Arogeti, F. L. Lewis, Static output-feedback h_∞ control design procedures for continuous-time systems with different levels of model knowledge, *IEEE Transactions on Cybernetics* (2021).
- [29] F. L. Lewis, D. Vrabie, V. L. Syrmos, *Optimal control*, John Wiley & Sons, 2012.
- [30] J. Gadewadikar, F. L. Lewis, L. Xie, V. Kucera, M. Abu-Khalaf, Parameterization of all stabilizing h-infinity static state-feedback gains: application to output-feedback design, *Automatica* 43 (9) (2007) 1597–1604.
- [31] W. M. Haddad, V. Chellaboina, *Nonlinear dynamical systems and control*, Princeton university press, 2011.

- [32] B. M. Chen, Robust and H-infinity Control, Springer Science & Business Media, 2013.
- [33] P. Apkarian, D. Noll, A. Rondepierre, Mixed h_2 h-infinity control via nonsmooth optimization, SIAM Journal on Control and Optimization 47 (3) (2008) 1516–1546.
- [34] T. Basar, P. Bernhard, H-infinity optimal control and related minimax design problems: a dynamic game approach, Springer Science & Business Media, 2008.
- [35] D. Kleinman, On an iterative technique for riccati equation computations, IEEE Transactions on Automatic Control 13 (1) (1968) 114–115.
- [36] H. W. Knobloch, A. Isidori, D. Flockerzi, Topics in control theory, Vol. 22, Birkhäuser, 2012.
- [37] D. Moerder, A. Calise, Convergence of a numerical algorithm for calculating optimal output feedback gains, IEEE Transactions on Automatic Control 30 (9) (1985) 900–903.
- [38] B. Datta, Numerical methods for linear control systems, Vol. 1, Academic Press, 2004.
- [39] D. Vrabie, O. Pastravanu, M. Abu-Khalaf, F. L. Lewis, Adaptive optimal control for continuous-time linear systems based on policy iteration, Automatica 45 (2) (2009) 477–484.
- [40] C. Li, G. G. Yin, L. Guo, C.-Z. Xu, et al., State observability and observers of linear-time-invariant systems under irregular sampling and sensor limitations, IEEE Transactions on Automatic Control 56 (11) (2011) 2639–2654.
- [41] B. L. Stevens, F. L. Lewis, E. N. Johnson, Aircraft control and simulation: dynamics, controls design, and autonomous systems, John Wiley & Sons, 2015.