

ADVERSARIAL PHYSICS-INFORMED MACHINE LEARNING FOR ROBUST OPTIMAL SAFE PREDEFINED-TIME STABILIZATION: A GAME-THEORETIC APPROACH *

NICK-MARIOS T. KOKOLAKIS^{†§}, SHANQING LIU[‡], JÉRÔME DARBON[‡], RAHUL MANGHARAM[†], AND GEORGE EM KARNIADAKIS[‡]

Abstract. We develop a game-theoretic framework for adversarially robust optimal safe predefined-time stabilization of parameter-dependent nonlinear dynamical systems with nonquadratic cost functionals. Our approach ensures that all system trajectories remain within a specified admissible set and converge to equilibrium in a predefined time despite adversarial disturbances. The control problem is formulated as a two-player zero-sum differential game, where the controller is a minimizing player and the adversary a maximizing player. We derive sufficient conditions for the existence of a saddle-point solution and safe predefined-time stability using a barrier Lyapunov function that satisfies a differential inequality and the steady-state Hamilton-Jacobi-Isaacs (HJI) equation. To address the analytical intractability of solving the HJI equation, we introduce a physics-informed learning algorithm that robustly learns the Nash safely predefined-time stabilizing control strategy. Simulation results demonstrate the efficacy and resilience of the proposed method in ensuring robust optimal safe predefined-time stabilization under adversarial disturbances.

Key words. Safety-critical control, predefined-time stability, robust optimal feedback control, differential games, physics-informed neural networks.

AMS subject classifications. 68T07, 49L12, 93D40.

1. Introduction. In control systems engineering, *autonomy* pertains to controlled systems that can operate without human intervention. Systems with this capability are termed autonomous systems (ASs), including, but not limited to, self-driving cars, humanoid robots, and unmanned aerial vehicles [48]. However, numerous incidents of unintended AS crashes highlight that ASs are *safety-critical systems*. Hence, guaranteeing safety is crucial, yielding the emergence of *safe autonomy* [51]. The control systems community can enable safe autonomy by harnessing the advantages of *nonlinear control theory* [17], *optimal control theory* [34], *robust control theory* [4], and *game theory* [5] to equip ASs with control architectures that ensure safety, stability, robustness, and performance. However, enabling safe autonomy becomes considerably more challenging in the presence of adversarial disturbances, which may arise from cyber-physical attacks, and hence necessitating the design of strategic decision-making mechanisms that synthesize adversarially robust optimal safe policies guaranteeing disturbance rejection in a predefined time rather than in an infinite or finite time. By incorporating *predefined-time adversarial robustness* into the safe control design, ASs can achieve enhanced resilience and reliability in real-world applications.

Stability theory concerns the behavior of the system trajectories of a dynamical system when the system initial state is near an equilibrium state [17]. The notion of

*Submitted to the editors November 20, 2025.

Funding: The work of N.-M.T.K and R.M. was partially supported by US DoT Safety21 National University Transportation Center and NSF under Grant CISE-2431569. The work of S.L., J.D., and G.E.K. was partially supported by Laboratory University Collaboration Initiative (LUCI) sponsored by the Basic Research Office, Office of Under Secretary of Defense for Research and Engineering (OUSD R&E).

[†]Department of Electrical and Systems Engineering, University of Pennsylvania, Philadelphia, PA 19104 USA ({nmkoko, rahulm}@seas.upenn.edu).

[‡]Division of Applied Mathematics, Brown University, Providence, RI 02912 USA ({shanqing.liu, jerome.darbon, george.karniadakis}@brown.edu).

[§]Corresponding author.

asymptotic stability in dynamical systems allows the convergence of system trajectories to a Lyapunov stable equilibrium point over the infinite horizon [17]. In contrast, the concept of *finite-time stability* enables the convergence of the system solutions to a Lyapunov stable equilibrium state in finite-time [10]. An inherent drawback of finite-time stability is that the settling-time function is not uniformly bounded, and hence the time of convergence to the equilibrium point may increase (possibly unboundedly) as the vector norm of the initial condition increases. The stronger notion of *fixed-time stability* ensures the convergence of the system trajectories to a finite-time stable equilibrium point in a fixed-time involving a uniformly bounded settling-time function [41]. A key limitation of fixed-time stability is that the upper bound of the settling-time function may not necessarily be predefined. Alternatively, the concept of *predefined-time stability* [46, 22] involves fixed-time stable *parameter-dependent* dynamical systems whose upper bound of the settling-time function can be chosen via an appropriate selection of the system parameters.

Optimal control theory concerns the synthesis of a control law for a given dynamical system to optimize a user-defined cost functional [34]. The notions of *optimality* and *stability* are intertwined in *optimal feedback stabilization* control problems that involve the synthesis of feedback control laws that guarantee closed-loop system stability while optimizing a given cost functional. The problems of optimal asymptotic, finite-time, fixed-time, and predefined-time stabilization are addressed in [6, 18, 27, 23]. Although these optimal control frameworks provide stability and performance guarantees, they do not account for adversarial disturbances, that is, strategic inputs designed to degrade performance or destabilize the system. Alternatively, *differential game theory* analyzes strategic interactions between competing agents, providing a powerful framework for addressing robust optimal control problems. Specifically, the disturbance rejection control problem in $\mathcal{H}_\infty$ control theory [3, 49, 4] can be formulated as a *two-player zero-sum game* wherein the controller is a minimizing player and the disturbance is a maximizing player. Infinite-horizon zero-sum games for linear and nonlinear dynamical systems with quadratic and nonquadratic cost functionals are studied in [20, 37, 35, 29, 30, 32, 31]. In particular, sufficient conditions are provided to characterize a saddle point solution for the game and closed-loop system asymptotic, partial-state asymptotic, and partial-state finite-time stability in [29, 30, 32]. In light of the above, although these game-theoretic control frameworks establish stability and Nash equilibrium, predefined-time stability is *not* ensured, and safety is *not* a design consideration.

Safe control theory concerns the controller analysis and synthesis for a safety critical dynamical system to guarantee the satisfaction of safety specifications [13], which can be expressed as forward invariance [17] of a set of safe system states. *Control barrier functions* (CBF) have been commonly used to guarantee the safety of a control system by rendering a safe set forward invariant [50, 2, 1]. The problem of *asymptotic* stabilization with guaranteed safety is addressed in [45] by merging a control Lyapunov function (CLF) [47] and a CBF [50]. However, a mapping that is both a CBF and a CLF cannot exist, as shown in [12]. Alternatively, quadratic programming (QP) has been utilized to combine a CLF and a CBF to synthesize controllers for safe *asymptotic* stabilization of nonlinear systems [44, 38]. A key limitation of these frameworks is the lack of robustness guarantees. In contrast, building on the results of [52], the notion of robust CBF is introduced in [21] to construct controllers for nonlinear systems with exogenous disturbances that guarantee safety and input-to-state stability [17]. However, CBF-based QPs introduce undesirable *asymptotically* stable equilibria [43], do not optimize closed-loop system performance in the face of a worst-

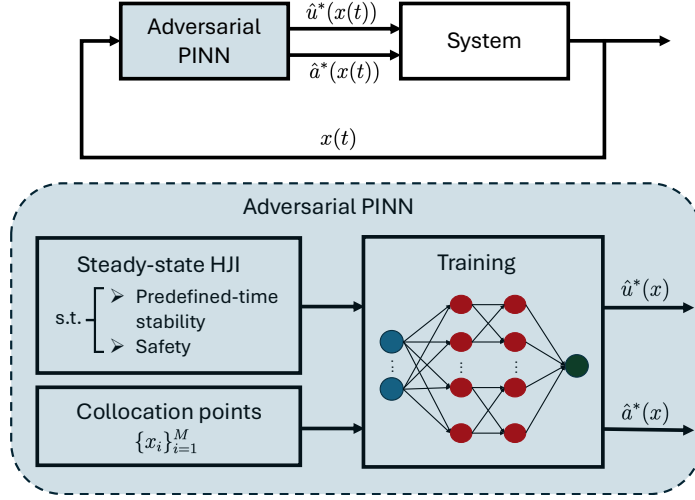


FIG. 1. Block diagram of the adversarial PINN control framework for robust optimal safe predefined-time stabilization. The adversarial PINN is trained offline using a set of collocation points to solve the steady-state HJI equation subject to predefined-time stability and safety constraints. The learned Nash control strategy $\hat{u}^*(\cdot)$ is then applied in feedback to ensure robust optimal safe predefined-time stabilization under the adversarial disturbance $\hat{a}^*(\cdot)$.

case adversary, and do not enforce time constraints. To the best of our knowledge, a game-theoretic control framework that *simultaneously* ensures *safety*, *predefined-time stability*, *optimality*, and *robustness* is absent from the literature.

To synthesize the Nash equilibrium strategies for a zero-sum differential game formulating a robust optimal stabilization problem, it is first necessary to determine the value of the game, which is a *stabilizing* solution of a nonlinear partial differential equation, the steady-state Hamilton–Jacobi–Isaacs (HJI) equation. Solving the steady-state HJI equation is challenging, except for special cases [33]. *Physics-informed neural networks* (PINNs), first developed in [42], have demonstrated their effectiveness in approximating a stabilizing solution to the steady-state Hamilton–Jacobi–Bellman (HJB) equation related to an *optimal stabilization problem* [16, 14, 28]. However, using PINNs to solve a zero-sum differential game has *not* been explored. To the best of our knowledge, the literature lacks a game-theoretic physics-informed learning framework that approximates the unique stabilizing solution to the steady-state HJI.

Contributions. The contributions of this paper are fourfold. First, we address an adversarially robust optimal safe predefined-time stabilization problem formulated as a two-player zero-sum differential game. Second, sufficient conditions for the existence of a saddle-point solution to the zero-sum game and closed-loop system safe predefined time stability are derived in terms of a barrier Lyapunov function. Third, a robust *inverse* optimal safe predefined time stabilization problem is addressed. Finally, an *adversarial physics-informed learning* algorithm is developed to learn the solution to the robust optimal safe predefined-time stabilization problem. A block diagram of the proposed adversarial PINN control framework is shown in Fig 1.

Structure. The remainder of the paper is structured as follows. The notion of safe predefined-time stability for general parameter-dependent nonlinear dynamical systems is introduced in Section 2, while Section 3 states the robust optimal safe

predefined-time stabilization problem. Section 4 introduces the robust optimal and inverse optimal safe predefined-time stabilization problem for parameter-dependent nonlinear *affine* dynamical systems. In Section 5, an adversarial physics-informed learning algorithm is developed to learn the solution to the robust optimal safe predefined-time stabilization problem. Section 6 presents illustrative numerical examples. Finally, conclusions are provided in Section 7.

Notation. The notation used in this paper is standard. Specifically, $\|\cdot\|_p \triangleq [\sum_{i=1}^n |x_i|^p]^{1/p}$, $1 \leq p < \infty$, denotes the ℓ^p -norm of a vector. We interchangeably use the notation $V'(x)$ and $V_x(x)$ to denote the gradient of a scalar-valued function $V(x)$ with respect to a vector-valued variable x , which is defined as a row vector. The signum function $\text{sgn} : \mathbb{R} \rightarrow \{-1, 0, 1\}$ is defined as $\text{sgn}(x) \triangleq x/|x|$, $x \neq 0$, and $\text{sgn}(0) \triangleq 0$. We define the gamma function $\Gamma(\cdot)$ as $\Gamma(x) \triangleq \int_0^\infty e^{-t} t^{x-1} dt$, $x > 0$. The indicator function of a set $A \subseteq \mathbb{R}^n$ is the function $\mathbb{1}_A : \mathbb{R}^n \rightarrow \{0, 1\}$ defined by $\mathbb{1}_A(x) \triangleq 1$, $x \in A$, and $\mathbb{1}_A(x) \triangleq 0$, $x \notin A$. Let $[\cdot]^\eta \triangleq |\cdot|^\eta \text{sgn}(\cdot)$, where $|\cdot|$ and $\text{sgn}(\cdot)$ operate component-wise and $\eta > 0$. The distance of a point $x_0 \in \mathbb{R}^n$ to a closed set $C \subseteq \mathbb{R}^n$ in the norm $\|\cdot\|$ is defined as $\text{dist}(x_0, C) \triangleq \inf_{x \in C} \{\|x_0 - x\|\}$. Finally, the notation $\partial\mathcal{S}$ denotes the boundary of the set $\mathcal{S}$.

2. Safe Predefined-Time Stability. In this section, we introduce the concept of *safe predefined-time stability* to characterize a class of nonlinear dynamical systems with the property that every trajectory starting in a given set of admissible states containing an equilibrium point remains in the set of admissible states and converges to the equilibrium point in a predefined time. Furthermore, we give sufficient conditions for safe predefined-time stability in terms of a barrier Lyapunov function.

To introduce the notions of *finite-* and *fixed-time stability*, consider the *parameter-independent* nonlinear dynamical system given by

$$(2.1) \quad \dot{x}(t) = f(x(t)), \quad x(0) = x_0, \quad t \geq 0,$$

where $x(t) \in \mathcal{D} \subseteq \mathbb{R}^n$, $t \geq 0$, is a system state vector, $\mathcal{D}$ is an open set with $0 \in \mathcal{D}$, $f : \mathcal{D} \rightarrow \mathbb{R}^n$ is continuous on $\mathcal{D}$ and satisfies $f(0) = 0$.

DEFINITION 2.1 ([10]). *The zero solution $x(t) \equiv 0$ to (2.1) is finite-time stable if it is Lyapunov stable and finite-time convergent, i.e., for all $x(0) \in \mathcal{N} \setminus \{0\}$, where $\mathcal{N} \subseteq \mathcal{D}$ is an open neighborhood of the origin, $\lim_{t \rightarrow T(x(0))} x(t) = 0$, where $T(\cdot)$ is the settling-time function such that $T(x(0)) < \infty$, $x(0) \in \mathcal{N}$. The zero solution $x(t) \equiv 0$ to (2.1) is globally finite-time stable if it is finite-time stable with $\mathcal{N} = \mathcal{D} = \mathbb{R}^n$. $\square$*

DEFINITION 2.2 ([41]). *The zero solution $x(t) \equiv 0$ to (2.1) is fixed-time stable if it is finite-time stable and the settling-time function $T(\cdot)$ is uniformly bounded, i.e., there exists $T_{\max} > 0$ such that $T(x(0)) \leq T_{\max}$, $x(0) \in \mathcal{N}$. The zero solution $x(t) \equiv 0$ to (2.1) is globally fixed-time stable if it is fixed-time stable with $\mathcal{N} = \mathcal{D} = \mathbb{R}^n$. $\square$*

The main difference between finite- and fixed-time stability is that in finite-time stability the settling-time function is not necessarily uniformly bounded, whereas fixed-time stability involves a uniformly bounded settling-time function.

Alternatively, to introduce the concept of *predefined-time stability*, consider the *parameter-dependent* nonlinear dynamical system given by

$$(2.2) \quad \dot{x}(t) = f(x(t), \theta), \quad x(0) = x_0, \quad t \geq 0,$$

where, for every $t \geq 0$, $x(t) \in \mathcal{D} \subseteq \mathbb{R}^n$ is a system state vector, $\mathcal{D}$ is an open set with $0 \in \mathcal{D}$, $\theta \in \mathbb{R}^N$ is a system parameter vector, $f : \mathcal{D} \times \mathbb{R}^N \rightarrow \mathbb{R}^n$ is such that $f(\cdot, \theta)$ is

continuous on $\mathcal{D}$ for all $\theta \in \mathbb{R}^N$ and $f(0, \cdot) = 0$. We write $s(t, x_0, \theta)$, $t \geq 0$, to denote the solution to (2.2) with initial condition x_0 and system parameter θ .

DEFINITION 2.3 ([22]). *The zero solution $x(t) \equiv 0$ to (2.2) is predefined-time stable with a predefined time $T_p > 0$ if there exists a system parameter vector $\theta \in \mathbb{R}^N$ such that the zero solution $x(t) \equiv 0$ to (2.2) is fixed-time stable with the settling-time function $T(\cdot, \theta)$ being uniformly bounded by T_p , i.e., $T(x(0), \theta) \leq T_p$, $x(0) \in \mathcal{N}$. The zero solution $x(t) \equiv 0$ to (2.2) is globally predefined-time stable with a predefined time $T_p > 0$ if it is predefined stable with a predefined time $T_p > 0$ with $\mathcal{N} = \mathcal{D} = \mathbb{R}^n$. $\square$*

REMARK 1. *Note that the system parameter vector θ implicitly depends on the predefined time T_p . $\square$*

The key difference between fixed- and predefined-time stability is that in predefined-time stability the upper bound of the settling-time function is defined a priori as an explicit function of the system parameters, whereas in fixed-time stability it is not necessarily predefined. It is important to note that fixed-time stability concerns the stability of a parameter-independent nonlinear system, unlike predefined-time stability, which involves parameter-dependent nonlinear systems. Lyapunov theorems for finite-, fixed-, and predefined-time stability are provided in [10, 41, 22].

A set of *admissible states* $\mathcal{S} \subset \mathcal{D}$ is called a *safe set* with respect to the parameter-dependent nonlinear dynamical system (2.2) if there exists a system parameter $\theta \in \mathbb{R}^N$ such that, for every $x_0 \in \mathcal{S}$, the solution $s(t, x_0, \theta)$, $t \geq 0$, to (2.2) satisfies $\text{dist}(s(t, x_0, \theta), \mathcal{S}) \equiv 0$. Equivalently, the set $\mathcal{S}$ is safe if there exists a system parameter vector such that $\mathcal{S}$ is positively invariant with respect to (2.2).

The next definition introduces the concept of *safe predefined-time stability*.

DEFINITION 2.4 ([26]). *Let $\mathcal{S} \subset \mathcal{D}$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. The zero solution $x(t) \equiv 0$ to (2.2) is safely predefined-time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$ if there exists a system parameter $\theta \in \mathbb{R}^N$ such that, for every $x_0 \in \mathcal{S}$, the solution $s(t, x_0, \theta)$, $t \geq 0$, to (2.2) satisfies $\text{dist}(s(t, x_0, \theta), \mathcal{S}) \equiv 0$ and $s(t, x_0, \theta) = 0$ for all $t \geq T_p$. $\square$*

REMARK 2. *Note that safe predefined-time stability unifies safety and predefined-time stability since positive invariance of the set of admissible states and predefined-time stability of the origin are established. Furthermore, note that the system parameter vector θ implicitly depends on the predefined time T_p and the set of admissible states $\mathcal{S}$. $\square$*

To present the main result of this section, the key definition of a *coercive* function is needed.

DEFINITION 2.5 ([8]). *Let $\mathcal{S} \subseteq \mathbb{R}^n$ be an unbounded set. A function $V : \mathcal{S} \rightarrow \mathbb{R}$ is called coercive if, for every sequence $\{x_k\}_{k=1}^{\infty}$ in $\mathcal{S}$ such that $\lim_{k \rightarrow \infty} \|x_k\| = \infty$, we have $\lim_{k \rightarrow \infty} V(x_k) = \infty$. $\square$*

The following key theorem gives sufficient conditions for safe predefined time stability of a parameter-dependent nonlinear dynamical system.

THEOREM 2.6 ([28]). *Consider the parameter-dependent nonlinear dynamical system (2.2). Let $\mathcal{S} \subset \mathcal{D}$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. Assume that there exist a continuously differentiable function $V : \mathcal{S} \rightarrow \mathbb{R}$, a system parameter vector $\theta \in \mathbb{R}^N$, and real numbers α , β , p , q , $k > 0$ such that*

$pk < 1$, $qk > 1$, and

$$(2.3) \quad V(0) = 0,$$

$$(2.4) \quad V(x) > 0, \quad x \in \mathcal{S} \setminus \{0\},$$

$$(2.5) \quad V(x) \rightarrow \infty \text{ as } x \rightarrow \partial\mathcal{S},$$

$$(2.6) \quad V'(x)f(x, \theta) \leq -\frac{\gamma}{T_p} \left(\alpha V^p(x) + \beta V^q(x) \right)^r, \quad x \in \mathcal{S},$$

where

$$\gamma \triangleq \frac{\Gamma\left(\frac{1-kp}{q-p}\right) \Gamma\left(\frac{kq-1}{q-p}\right)}{\alpha^k \Gamma(k)(q-p)} \left(\frac{\alpha}{\beta}\right)^{\frac{1-kp}{q-p}}.$$

If either $\mathcal{S}$ is bounded or both $\mathcal{S}$ is unbounded and $V(\cdot)$ is coercive, then the zero solution $x(t) \equiv 0$ to (2.2) is safely predefined time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$.

A continuously differentiable function $V(\cdot)$ satisfying (2.3)-(2.5) is called a *safely predefined-time stabilizing Lyapunov function candidate* for the nonlinear dynamical system (2.2). If, additionally, $V(\cdot)$ satisfies (2.6), $V(\cdot)$ is called a *safely predefined-time stabilizing Lyapunov function* for the nonlinear dynamical system (2.2).

REMARK 3. Theorem 2.6 investigates safe predefined-time stability for the nonlinear dynamical system (2.2) without requiring knowledge of the system trajectories. $\square$

3. Robust Optimal Safe Predefined-Time Stabilization. In this section, we formulate the robust optimal safe predefined time stabilization problem as a *two-player zero-sum differential game* to characterize robust optimal feedback controllers that render the equilibrium point of the closed-loop system safely predefined time while optimizing the closed-loop system performance against the worst-case feedback adversary. Specifically, we provide sufficient conditions for the existence of a safely predefined time stabilizing feedback *Nash equilibrium* solution to the zero-sum game.

Consider the controlled parameter-dependent nonlinear dynamical system given by

$$(3.1) \quad \dot{x}(t) = F(x(t), \theta_s, u(t), a(t)), \quad x(0) = x_0, \quad t \geq 0,$$

where, for every $t \geq 0$, $x(t) \in \mathcal{D} \subseteq \mathbb{R}^n$ is the state vector, $\mathcal{D}$ is an open set with $0 \in \mathcal{D}$, $\theta_s \in \mathbb{R}^{N_s}$ is the system parameter vector, $u(t) \in U \subseteq \mathbb{R}^{m_u}$ is the control input with $0 \in U$, $a(t) \in A \subseteq \mathbb{R}^{m_a}$ is the adversary input with $0 \in A$, and $F : \mathcal{D} \times \mathbb{R}^{N_s} \times U \times A \rightarrow \mathbb{R}^n$ is such that $F(\cdot, \theta_s, \cdot, \cdot)$ is jointly continuous on $\mathcal{D} \times U \times A$ for all $\theta_s \in \mathbb{R}^{N_s}$ and $F(0, \cdot, 0, 0) = 0$. The control $u(\cdot)$ and adversary $a(\cdot)$ in (3.1) belong to the class of *admissible controls* $\mathcal{U} \triangleq \{u : [0, \infty) \rightarrow U : u(\cdot) \text{ is Lebesgue measurable}\}$ and *admissible adversaries* $\mathcal{A} \triangleq \{a : [0, \infty) \rightarrow A : a(\cdot) \text{ is Lebesgue measurable}\}$. We assume that the required properties for the existence and uniqueness of solutions to (3.1) are satisfied, and we write $s(t, x_0, \theta_s, u(\cdot), a(\cdot))$, $t \geq 0$, to denote the solution to (3.1) with initial condition x_0 , system parameter θ_s , admissible control $u(\cdot)$, and admissible adversary $a(\cdot)$.

The mappings $u^* : \mathcal{S} \times \mathbb{R}^{N_c} \rightarrow U$ and $a^* : \mathcal{S} \times \mathbb{R}^{N_a} \rightarrow A$ such that $u^*(\cdot, \theta_c)$ and $a^*(\cdot, \theta_a)$ are Lebesgue measurable functions on $\mathcal{S}$ for every control parameter vector $\theta_c \in \mathbb{R}^{N_c}$ and adversary parameter vector $\theta_a \in \mathbb{R}^{N_a}$ with $u^*(0, \cdot) = 0$ and

$a^*(0, \cdot) = 0$ are called a *control strategy* and an *adversary strategy*. Furthermore, if $u(t) = u^*(x(t), \theta_c)$, $t \geq 0$, and $a(t) = a^*(x(t), \theta_a)$, $t \geq 0$, where $u^*(\cdot, \cdot)$ is a control strategy and $a^*(\cdot, \cdot)$ is an adversary strategy and $x(t)$ is a solution to (3.1), then we call $u(\cdot)$ a *feedback control strategy* and $a(\cdot)$ a *feedback adversary strategy*. Note that a feedback control strategy is an admissible control and a feedback adversary strategy is an admissible adversary since $u^*(\cdot, \theta_c)$ and $a^*(\cdot, \theta_a)$ are Lebesgue measurable functions and take values in U and A .

Given a control strategy $u^*(\cdot, \cdot)$ and an adversary strategy $a^*(\cdot, \cdot)$ and a feedback control strategy $u(t) = u^*(x(t), \theta_c)$, $t \geq 0$, and a feedback adversary strategy $a(t) = a^*(x(t), \theta_a)$, $t \geq 0$, the closed-loop system is given by

$$(3.2) \quad \dot{x}(t) = F(x(t), \theta_s, u^*(x(t), \theta_c), a^*(x(t), \theta_a)), \quad x(0) = x_0, \quad t \geq 0,$$

which can be cast in the form of (2.2) with $N = N_s + N_c + N_a$, $\theta = [\theta_s^T, \theta_c^T, \theta_a^T]^T$, and $f(x, \theta) = F(x, \theta_s, u^*(x, \theta_c), a^*(x, \theta_a))$.

We now define the notion of a *safely predefined-time stabilizing pair of feedback strategies*.

DEFINITION 3.1. *Consider the controlled parameter-dependent nonlinear dynamical system given by (3.1). Let $\mathcal{S} \subset \mathcal{D}$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. The pair $(u^*(\cdot), a^*(\cdot))$ of the feedback control strategy $u(\cdot) = u^*(x(\cdot), \theta_c)$ and the feedback adversary strategy $a(\cdot) = a^*(x(\cdot), \theta_a)$ is safely predefined-time stabilizing if there exists a system parameter $\theta_s \in \mathbb{R}^{N_s}$ such that the zero solution $x(t) \equiv 0$ to the closed-loop system (3.2) is safely predefined-time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$. $\square$*

Given a set of admissible states $\mathcal{S} \subset \mathcal{D}$ with $0 \in \mathcal{S}$ and a predefined time $T_p > 0$, we define, for every $x_0 \in \mathcal{S}$, the set of *safely predefined-time stabilizing pairs* of feedback strategies by $\mathcal{F}(x_0, \mathcal{S}, T_p) \triangleq \{u : [0, \infty) \rightarrow U \text{ and } a : [0, \infty) \rightarrow A : u(\cdot) \text{ and } a(\cdot) \text{ are a feedback control strategy and a feedback adversary strategy and } s(t, x_0, \theta_s, u(\cdot), a(\cdot)), t \geq 0, \text{ is a solution to (3.1) satisfying } \text{dist}(s(t, x_0, \theta_s, u(\cdot), a(\cdot)), \mathcal{S}) \equiv 0 \text{ and } s(t, x_0, \theta_s, u(\cdot), a(\cdot)) = 0 \text{ for all } t \geq T_p\} \subset \mathcal{U} \times \mathcal{A}$. Furthermore, let $\mathcal{F}_u(x_0, \mathcal{S}, T_p) \subset \mathcal{U}$ and $\mathcal{F}_a(x_0, \mathcal{S}, T_p) \subset \mathcal{A}$ be the set of feedback control strategies and feedback adversary strategies such that $\mathcal{F}_u(x_0, \mathcal{S}, T_p) \times \mathcal{F}_a(x_0, \mathcal{S}, T_p) = \mathcal{F}(x_0, \mathcal{S}, T_p)$.

To evaluate the performance of the controlled parameter-dependent nonlinear dynamical system (3.1) over the time interval $[0, T_p]$ for a given predefined time $T_p > 0$ and a set of admissible states $\mathcal{S}$ with $0 \in \mathcal{S}$, we define, for every $x_0 \in \mathcal{S}$, $\theta_s \in \mathbb{R}^{N_s}$, $u(\cdot) \in \mathcal{U}$, and $a(\cdot) \in \mathcal{A}$, the cost functional

$$(3.3) \quad J(x_0, \theta_s, u(\cdot), a(\cdot)) \triangleq \int_0^{T_p} r(x(t), u(t), a(t)) dt,$$

where $r : \mathcal{S} \times U \times A \rightarrow \mathbb{R}$ is jointly continuous in x , u , and a . The controller $u(\cdot)$ is a player that *minimizes* the cost functional (3.3), while the adversary $a(\cdot)$ is a player that *maximizes* the cost functional (3.3).

In light of the above, we now state the robust optimal safe predefined time stabilization problem as a *two-player zero-sum game*.

PROBLEM 1. *Consider the controlled parameter-dependent nonlinear dynamical system given by (3.1) with the cost functional (3.3). Let the controller $u(\cdot)$ be a player that minimizes the cost functional (3.3) and let the adversary $a(\cdot)$ be a player that maximizes the cost functional (3.3). Let $\mathcal{S} \subset \mathcal{D}$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. For every $x_0 \in \mathcal{S}$, let $\mathcal{F}(x_0, \mathcal{S}, T_p) \subset$*

$\mathcal{U} \times \mathcal{A}$ be the set of safely predefined time stabilizing pairs of feedback strategies and suppose that $\mathcal{F}(x_0, \mathcal{S}, T_p)$ is nonempty. For every initial condition $x_0 \in \mathcal{S}$, synthesize $(u^*(\cdot), a^*(\cdot)) \in \mathcal{F}(x_0, \mathcal{S}, T_p)$ such that the zero solution $x(t) \equiv 0$ to the closed-loop system (3.2) is safely predefined-time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$ and $(u^*(\cdot), a^*(\cdot))$ is a saddle point of the cost functional (3.3). $\square$

In other words, the robust optimal safe predefined-time stabilization problem involves synthesizing a feedback control strategy $u^*(\cdot)$ to ensure safe predefined-time stability while optimizing performance against the worst-case feedback adversary strategy $a^*(\cdot)$ so that $(u^*(\cdot), a^*(\cdot))$ is a saddle point of the cost functional (3.3).

For every $x_0 \in \mathcal{S}$ and $\theta_s \in \mathbb{R}^{N_s}$, the robust optimal safe predefined-time stabilization problem involves the minimax control problem

$$\min_{u(\cdot) \in \mathcal{F}_u(x_0, \mathcal{S}, T_p)} \max_{a(\cdot) \in \mathcal{F}_a(x_0, \mathcal{S}, T_p)} J(x_0, \theta_s, u(\cdot), a(\cdot))$$

subject to (3.1).

The next theorem gives sufficient conditions for the existence of a safely predefined time stabilizing feedback *Nash equilibrium* solution to the two-player zero-sum game. For the statement of this result, we need to define the Hamiltonian function

$$(3.4) \quad \begin{aligned} H(x, \theta_s, u, a, \lambda) &\triangleq r(x, u, a) + \lambda^T F(x, \theta_s, u, a), \\ (x, \theta_s, u, a, \lambda) &\in \mathcal{S} \times \mathbb{R}^{N_s} \times U \times A \times \mathbb{R}^n. \end{aligned}$$

THEOREM 3.2. *Consider the controlled parameter-dependent nonlinear dynamical system given by (3.1) with cost functional (3.3). Let the controller $u(\cdot)$ be a player that minimizes the cost functional (3.3) and let the adversary $a(\cdot)$ be a player that maximizes the cost functional (3.3). Let $\mathcal{S} \subset \mathcal{D}$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. For every $x_0 \in \mathcal{S}$, let $\mathcal{F}(x_0, \mathcal{S}, T_p) \subset \mathcal{U} \times \mathcal{A}$ be the set of safely predefined time stabilizing pairs of feedback strategies and suppose that $\mathcal{F}(x_0, \mathcal{S}, T_p)$ is nonempty. Assume that there exist a continuously differentiable function $V : \mathcal{S} \rightarrow \mathbb{R}$, a system parameter vector $\theta_s \in \mathbb{R}^{N_s}$, real numbers $\alpha, \beta, p, q, r > 0$ such that $pr < 1$ and $qr > 1$, a feedback control strategy $u^* : \mathcal{S} \times \mathbb{R}^{N_c} \rightarrow U$, and a feedback adversary strategy $a^* : \mathcal{S} \times \mathbb{R}^{N_a} \rightarrow A$ such that*

$$(3.5) \quad H(x, \theta_s, u^*(x, \theta_c), a^*(x, \theta_a), V'^T(x)) = 0, \quad x \in \mathcal{S},$$

$$(3.6) \quad H(x, \theta_s, u^*(x, \theta_c), a, V'^T(x)) \leq 0, \quad (x, a) \in \mathcal{S} \times A,$$

$$(3.7) \quad H(x, \theta_s, u, a^*(x, \theta_a), V'^T(x)) \geq 0, \quad (x, u) \in \mathcal{S} \times U,$$

$$(3.8) \quad u^*(0, \theta_c) = 0,$$

$$(3.9) \quad a^*(0, \theta_a) = 0,$$

$$(3.10) \quad V(0) = 0,$$

$$(3.11) \quad V(x) > 0, \quad x \in \mathcal{S} \setminus \{0\},$$

$$(3.12) \quad V(x) \rightarrow \infty \text{ as } x \rightarrow \partial\mathcal{S},$$

$$(3.13) \quad V'(x)F(x, \theta_s, u^*(x, \theta_c), a^*(x, \theta_a)) \leq -\frac{\gamma}{T_p} \left(\alpha V^p(x) + \beta V^q(x) \right)^r, \quad x \in \mathcal{S},$$

where

$$\gamma \triangleq \frac{\Gamma\left(\frac{1-rp}{q-p}\right)\Gamma\left(\frac{rq-1}{q-p}\right)}{\alpha^r\Gamma(r)(q-p)}\left(\frac{\alpha}{\beta}\right)^{\frac{1-rp}{q-p}}.$$

If either $\mathcal{S}$ is bounded or both $\mathcal{S}$ is unbounded and $V(\cdot)$ is coercive, then with the feedback control strategy $u(\cdot) = u^*(x(\cdot), \theta_c)$ and the feedback adversary strategy $a(\cdot) = a^*(x(\cdot), \theta_a)$, the zero solution $x(t) \equiv 0$ to (3.2) is safely predefined-time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$. Furthermore, if $x_0 \in \mathcal{S}$, then the pair of feedback strategies $(u^*(\cdot), d^*(\cdot))$ is the Nash equilibrium of the two-player zero-sum game in the sense that

$$\begin{aligned} J(x_0, \theta_s, u^*(x(\cdot), \theta_c), a^*(x(\cdot), \theta_a)) &= \min_{u(\cdot) \in \mathcal{F}_u(x_0, \mathcal{S}, T_p)} \max_{a(\cdot) \in \mathcal{F}_a(x_0, \mathcal{S}, T_p)} J(x_0, \theta_s, u(\cdot), a(\cdot)) \\ (3.14) \quad &= \max_{a(\cdot) \in \mathcal{F}_a(x_0, \mathcal{S}, T_p)} \min_{u(\cdot) \in \mathcal{F}_u(x_0, \mathcal{S}, T_p)} J(x_0, \theta_s, u(\cdot), a(\cdot)) \end{aligned}$$

and the Nash value is

$$(3.15) \quad J(x_0, \theta_s, u^*(x(\cdot), \theta_c), a^*(x(\cdot), \theta_a)) = V(x_0).$$

Proof. Safe predefined-time stability is a direct consequence of (3.10)–(3.12) and (3.13) by applying Theorem 2.6 to the closed-loop system (3.2).

Next, to prove Nash equilibrium, let $x_0 \in \mathcal{S}$, let $u(\cdot) = u^*(x(\cdot), \theta_c)$, let $a(\cdot) = a^*(x(\cdot), \theta_a)$, and let $x(t) = s(t, x_0, \theta_s, u^*(x(\cdot), \theta_c), a^*(x(\cdot), \theta_a))$, $t \geq 0$, be the solution to (3.2). Then, since the time derivative of $V(\cdot)$ along the trajectories of the closed-loop system (3.2) is defined by

$$\dot{V}(x(t)) \triangleq V'(x(t))F(x(t), \theta_s, u^*(x(t), \theta_c), a^*(x(t), \theta_a)), \quad t \geq 0,$$

and safe predefined-time stability of (3.2) implies $x(t) \in \mathcal{S}$, $t \geq 0$, it follows from (3.5) that

$$\begin{aligned} &H(x(t), \theta_s, u^*(x(t), \theta_c), a^*(x(t), \theta_a), V'^T(x(t))) \\ &= r(x, u^*(x(t), \theta_c), a^*(x(t), \theta_a)) + \dot{V}(x(t)) \\ (3.16) \quad &= 0, \quad x(t) \in \mathcal{S}, \quad t \geq 0. \end{aligned}$$

Now, integrating (3.16) over the time interval $[0, T_p]$ and using the definition of the cost functional (3.3) yields

$$J(x_0, \theta_s, u^*(x(\cdot), \theta_c), a^*(x(\cdot), \theta_a)) + V(x(T_p)) = V(x_0),$$

which, since $V(x(T_p)) = 0$ by (3.10) and safe predefined-time stability of (3.2), implies (3.15).

Next, let $x_0 \in \mathcal{S}$, let $u(\cdot) = u^*(x(\cdot), \theta_c)$, let $a(\cdot) \in \mathcal{F}_a(x_0, \mathcal{S}, T_p)$, and let $x(t) = s(t, x_0, \theta_s, u^*(x(\cdot), \theta_c), a(x(\cdot), \theta_a))$, $t \geq 0$, be the solution to

$$(3.17) \quad \dot{x}(t) = F(x(t), \theta_s, u^*(x(t), \theta_c), a(x(t), \theta_a)), \quad x(0) = x_0, \quad t \geq 0.$$

Then, since the time derivative of $V(\cdot)$ along the trajectories of (3.17) is given by

$$\dot{V}(x(t)) \triangleq V'(x(t))F(x(t), \theta_s, u^*(x(t), \theta_c), a(x(t), \theta_a)), \quad t \geq 0,$$

and $(u^*(\cdot), a(\cdot)) \in \mathcal{F}(x_0, \mathcal{S}, T_p)$ implies $x(t) \in \mathcal{S}$, $t \geq 0$, it follows from (3.6) that

$$\begin{aligned} H(x(t), \theta_s, u^*(x(t), \theta_c), d(x(t), \theta_a), V^T(x(t))) \\ = r(x(t), u^*(x(t), \theta_c), d(x(t), \theta_a)) + \dot{V}(x(t)) \\ \leq 0, \quad x(t) \in \mathcal{S}, \quad t \geq 0. \end{aligned} \quad (3.18)$$

Now, integrating (3.18) over $[0, T_p]$, using (3.3), (3.10), and (3.15), and using the fact that

$(u^*(\cdot), a(\cdot)) \in \mathcal{F}(x_0, \mathcal{S}, T_p)$ yields

$$\begin{aligned} J(x_0, \theta_s, u^*(\cdot), a(\cdot)) &\leq V(x_0) \\ &= J(x_0, \theta_s, u^*(\cdot), a^*(\cdot)). \end{aligned} \quad (3.19)$$

Next, let $x_0 \in \mathcal{S}$, let $u(\cdot) \in \mathcal{F}_u(x_0, \mathcal{S}, T_p)$, let $a(\cdot) = a^*(x(\cdot), \theta_a)$, and let $x(t) = s(t, x_0, \theta_s, u(x(\cdot), \theta_c), a^*(x(\cdot), \theta_a))$, $t \geq 0$, be the solution to

$$\dot{x}(t) = F(x(t), \theta_s, u(x(t), \theta_c), a^*(x(t), \theta_a)), \quad x(0) = x_0, \quad t \geq 0. \quad (3.20)$$

Then, computing the time derivative of $V(\cdot)$ along the trajectories of (3.20) as

$$\dot{V}(x(t)) \triangleq V'(x(t))F(x(t), \theta_s, u(x(t), \theta_c), a^*(x(t), \theta_a)), \quad t \geq 0,$$

and using the fact that $(u(\cdot), a^*(\cdot)) \in \mathcal{F}(x_0, \mathcal{S}, T_p)$ implies $x(t) \in \mathcal{S}$, $t \geq 0$, it follows from (3.7) that

$$\begin{aligned} H(x(t), \theta_s, u(x(t), \theta_c), a^*(x(t), \theta_a), V^T(x(t))) \\ = r(x(t), u(x(t), \theta_c), a^*(x(t), \theta_a)) + \dot{V}(x(t)) \\ \geq 0, \quad x(t) \in \mathcal{S}, \quad t \geq 0. \end{aligned} \quad (3.21)$$

Now, integrating (3.21) over $[0, T_p]$ and using (3.3), (3.10), (3.15), and the fact that $(u(\cdot), a^*(\cdot)) \in \mathcal{F}(x_0, \mathcal{S}, T_p)$, we obtain

$$\begin{aligned} J(x_0, \theta_s, u(\cdot), a^*(\cdot)) &\geq V(x_0) \\ &= J(x_0, \theta_s, u^*(\cdot), a^*(\cdot)). \end{aligned} \quad (3.22)$$

In light of (3.19) and (3.22), it follows that, for every $x_0 \in \mathcal{S}$, $u(\cdot) \in \mathcal{F}_u(x_0, \mathcal{S}, T_p)$, and $a(\cdot) \in \mathcal{F}_a(x_0, \mathcal{S}, T_p)$,

$$\begin{aligned} J(x_0, \theta_s, u^*(\cdot), a(\cdot)) &\leq J(x_0, \theta_s, u^*(\cdot), a^*(\cdot)) \\ &\leq J(x_0, \theta_s, u(\cdot), a^*(\cdot)), \end{aligned}$$

which implies (3.14). ■

It is important to note that, unlike QP-based methods that depend on system trajectories, Theorem 3.2 examines robust optimal safe predefined-time stabilization for the the controlled parameter-dependent nonlinear dynamical system (3.1) with the cost functional (3.3) without knowledge of the system trajectories.

REMARK 4. *The following observations provide insights into the robust optimal safe predefined-time stabilization problem.*

1) Equation (3.14) asserts that, for all $x_0 \in \mathcal{S}$, $(u^*(\cdot), d^*(\cdot))$ is a saddle point equilibrium for the two-player zero sum game over the set of safely predefined time

stabilizing pairs of feedback strategies $\mathcal{F}(x_0, \mathcal{S}, T_p)$, where no player benefits from a unilateral deviation from the equilibrium. Note that an explicit characterization of $\mathcal{F}_u(x_0, \mathcal{S}, T_p)$, $\mathcal{F}_a(x_0, \mathcal{S}, T_p)$, and $\mathcal{F}(x_0, \mathcal{S}, T_p)$ is not required.

2) Conditions (3.10)-(3.13) guarantee safe predefined time stability, while conditions (3.5)-(3.7) establish Nash equilibrium over the set of admissible control values U and adversary values A .

3) Equation (3.5) is the steady-state HJI equation for the controlled parameter-dependent nonlinear dynamical system (3.1) with the cost functional (3.3). Furthermore, note that (3.5), (3.6), and (3.7) imply

$$\min_{u \in U} \max_{a \in A} H(x, \theta_s, u, a, V'^T(x)) = \max_{a \in A} \min_{u \in U} H(x, \theta_s, u, a, V'^T(x)) = 0, \quad x \in \mathcal{S},$$

which is an alternative expression for the steady-state Hamilton-Jacobi-Isaacs equation (3.5).

4) The Nash feedback control strategy $u^*(x, \theta_c)$ is an adversarially robust optimal control that minimizes the maximum value of the Hamiltonian function (3.4) for every $x \in \mathcal{S}$ with respect to U , that is,

$$\begin{aligned} u^*(x, \theta_c) &= \arg \min_{u \in U} \max_{a \in A} H(x, \theta_s, u, a, V'^T(x)) \\ &= \arg \min_{u \in U} H(x, \theta_s, u, a^*(x, \theta_a), V'^T(x)). \end{aligned}$$

Hence, it follows that $u^*(\cdot, \theta_c)$ is the optimal response against the worst-case adversary $a^*(\cdot, \theta_a)$ regardless of the initial condition $x_0 \in \mathcal{S}$.

5) Equation (3.15) shows that the safely predefined time stabilizing Lyapunov function $V(\cdot)$ for the closed-loop system (3.2) serves as a value of the game quantifying the best worst-case performance or, equivalently, the worst best-case performance.

6) Although an explicit expression for the settling-time function $T(\cdot, \cdot)$ cannot be given, safe predefined time stability guarantees that there exists a parameter vector $\theta \in \mathbb{R}^N$ such that the settling-time function $T(\cdot, \theta)$ is uniformly bounded by T_p , that is, $\sup_{x_0 \in \mathcal{S}} T(x_0, \theta) \leq T_p$.

7) The system parameter vector $\theta_s \in \mathbb{R}^{N_s}$, the control parameter vector $\theta_c \in \mathbb{R}^{N_c}$, and the adversary parameter vector $\theta_a \in \mathbb{R}^{N_a}$ implicitly depend on the predefined time T_p and the set of admissible states $\mathcal{S}$. $\square$

4. Robust Optimal and Inverse Optimal Safe Predefined-Time Stabilization for Nonlinear Affine Systems. In this section, we specialize the results of Section 3 to parameter-dependent nonlinear affine dynamical systems of the form

$$(4.1) \quad \dot{x}(t) = f(x(t), \theta_f) + G(x(t), \theta_G)u(t) + K(x(t), \theta_K)a(t), \quad x(0) = x_0, \quad t \geq 0,$$

where, for every $t \geq 0$, $x(t) \in \mathbb{R}^n$ is the state vector, $u(t) \in \mathbb{R}^{m_u}$ is the control input, $a(t) \in \mathbb{R}^{m_a}$ is the adversary input, $f : \mathbb{R}^n \times \mathbb{R}^{N_f} \rightarrow \mathbb{R}^n$ is such that $f(\cdot, \theta_f)$ is continuous on $\mathbb{R}^n$ for all $\theta_f \in \mathbb{R}^{N_f}$ and $f(0, \cdot) = 0$, $G : \mathbb{R}^n \times \mathbb{R}^{N_G} \rightarrow \mathbb{R}^{n \times m_u}$ is such that $G(\cdot, \theta_G)$ is continuous on $\mathbb{R}^n$ for all $\theta_G \in \mathbb{R}^{N_G}$, and $K : \mathbb{R}^n \times \mathbb{R}^{N_K} \rightarrow \mathbb{R}^{n \times m_a}$ is such that $K(\cdot, \theta_K)$ is continuous on $\mathbb{R}^n$ for all $\theta_K \in \mathbb{R}^{N_K}$. In this case, the system parameter vector $\theta_s \in \mathbb{R}^{N_s}$ is defined as $\theta_s \triangleq [\theta_f^T, \theta_G^T, \theta_K^T]^T$.

For every $(x, u, a) \in \mathcal{S} \times \mathbb{R}^{m_u} \times \mathbb{R}^{m_a}$, we consider running costs $r(x, u, a)$ of the form

$$(4.2) \quad r(x, u, a) \triangleq L(x) + L_u(x)u + L_a(x)a + u^T R_u(x)u - a^T R_a(x)a,$$

where $L : \mathcal{S} \rightarrow \mathbb{R}$, $L_u : \mathcal{S} \rightarrow \mathbb{R}^{1 \times m_u}$, $L_a : \mathcal{S} \rightarrow \mathbb{R}^{1 \times m_a}$, $R_u : \mathcal{S} \rightarrow \mathbb{R}^{m_u \times m_u}$, and $R_a : \mathcal{S} \rightarrow \mathbb{R}^{m_a \times m_a}$ are continuous on $\mathcal{S}$ such that $L(0) = 0$, $L_u(0) = 0$, $L_a(0) = 0$, $R_u(x) > 0$, $x \in \mathcal{S}$, and $R_a(x) > 0$, $x \in \mathcal{S}$. Note that $L(\cdot)$ penalizes deviations of the state x from the origin, if $L(x) \geq 0$, $x \in \mathcal{S}$, $L_u(\cdot)$ penalizes alignment with the control input u , $L_a(\cdot)$ rewards alignment with the adversary input a , and $R_u(\cdot)$ and $R_a(\cdot)$ penalize control and adversary effort.

REMARK 5. The functions $L_u(x)$ and $L_a(x)$ in the running cost (4.2) are arbitrary vector functions of $x \in \mathcal{S}$ and provide the cross-weighting terms $L_u(x)u$ and $L_a(x)a$ in the running cost $r(x, u, a)$. Note that the terms $L_u(x)u$ and $L_a(x)a$ can be written as $L_u(x)u = \|L_u^T(x)\|_2 \|u\|_2 \cos(\phi_{L_u}(x) - \phi_u)$, $x \in \mathcal{S}$, and $L_a(x)a = \|L_a^T(x)\|_2 \|a\|_2 \cos(\phi_{L_a}(x) - \phi_a)$, $x \in \mathcal{S}$, where $\phi_{L_u}(x) - \phi_u$ is the angular difference between $L_u^T(x)$ and u at $x \in \mathcal{S}$ and $\phi_{L_a}(x) - \phi_a$ is the angular difference between $L_a^T(x)$ and a at $x \in \mathcal{S}$. Hence, $L_u(x)$ penalizes alignment with the control input u at $x \in \mathcal{S}$ and $L_a(x)$ rewards alignment with the adversary input a at $x \in \mathcal{S}$. For example, $L_u(\cdot)$ and $L_a(\cdot)$ may be designed to indicate the control and adversary input directions that yield instability, unsafety, or low performance. $\square$

For every $x_0 \in \mathcal{S}$, $\theta_s \in \mathbb{R}^{N_s}$, $u(\cdot) \in \mathcal{U}$, and $a(\cdot) \in \mathcal{A}$, the cost functional (3.3), with $r(\cdot, \cdot, \cdot)$ given by (4.2), becomes

$$(4.3) \quad J(x_0, \theta_s, u(\cdot), a(\cdot)) = \int_0^{T_p} \left(L(x(t)) + L_u(x(t))u(t) + L_a(x(t))a(t) + u^T(t)R_u(x(t))u(t) - a^T(t)R_a(x(t))a(t) \right) dt.$$

Next, we specialize Theorem 3.2 to parameter-dependent nonlinear affine dynamical systems (4.1) with the cost functional (4.3).

COROLLARY 4.1. Consider the parameter-dependent nonlinear affine dynamical system (4.1) with cost functional (3.3). Let the controller $u(\cdot)$ be a player that minimizes the cost functional (3.3) and let the adversary $a(\cdot)$ be a player that maximizes the cost functional (3.3). Let $\mathcal{S} \subset \mathbb{R}^n$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. For every $x_0 \in \mathcal{S}$, let $\mathcal{F}(x_0, \mathcal{S}, T_p) \subset \mathcal{U} \times \mathcal{D}$ be the set of safely predefined time stabilizing pairs of feedback strategies and suppose that $\mathcal{F}(x_0, \mathcal{S}, T_p)$ is nonempty. Assume that there exist a continuously differentiable function $V : \mathcal{S} \rightarrow \mathbb{R}$, system parameter vectors $\theta_f \in \mathbb{R}^{N_f}$, $\theta_G \in \mathbb{R}^{N_G}$, and $\theta_K \in \mathbb{R}^{N_K}$, and

real numbers $\alpha, \beta, p, q, r > 0$ such that $pr < 1$, $qr > 1$, and

$$(4.4) \quad \begin{aligned} & L(x) + V'(x)f(x, \theta_f) - \frac{1}{4} \left[V'(x)G(x, \theta_G) + L_u(x) \right] \\ & \cdot R_u^{-1}(x) \left[V'(x)G(x, \theta_G) + L_u(x) \right]^T + \frac{1}{4} \left[V'(x)K(x, \theta_K) + L_a(x) \right] \\ & \cdot R_a^{-1}(x) \left[V'(x)K(x, \theta_K) + L_a(x) \right]^T = 0, \quad x \in \mathcal{S}, \end{aligned}$$

$$(4.5) \quad V(0) = 0,$$

$$(4.6) \quad V(x) > 0, \quad x \in \mathcal{S} \setminus \{0\},$$

$$(4.7) \quad V(x) \rightarrow \infty \text{ as } x \rightarrow \partial\mathcal{S},$$

$$(4.8) \quad \begin{aligned} & V'(x) \left[f(x, \theta_f) - \frac{1}{2} G(x, \theta_G) R_u^{-1}(x) L_u^T(x) + \frac{1}{2} K(x, \theta_K) R_a^{-1}(x) L_a^T(x) \right. \\ & \left. - \frac{1}{2} G(x, \theta_G) R_u^{-1}(x) G^T(x, \theta_G) V'^T(x) + \frac{1}{2} K(x, \theta_K) R_a^{-1}(x) K^T(x, \theta_K) V'^T(x) \right] \\ & \leq -\frac{\gamma}{T_p} (\alpha V^p(x) + \beta V^q(x))^r, \quad x \in \mathcal{S}, \end{aligned}$$

where

$$\gamma \triangleq \frac{\Gamma\left(\frac{1-rp}{q-p}\right) \Gamma\left(\frac{rq-1}{q-p}\right)}{\alpha^r \Gamma(r)(q-p)} \left(\frac{\alpha}{\beta}\right)^{\frac{1-rp}{q-p}}.$$

If either $\mathcal{S}$ is bounded or both $\mathcal{S}$ is unbounded and $V(\cdot)$ is coercive, then there exist a control parameter vector $\theta_c \in \mathbb{R}^{N_c}$ and an adversary parameter vector $\theta_a \in \mathbb{R}^{N_a}$ such that with the feedback control strategy

$$(4.9) \quad \begin{aligned} & u(t) = u^*(x(t), \theta_c) \\ & \triangleq -\frac{1}{2} R_u^{-1}(x(t)) [L_u(x(t)) + V'(x(t))G(x(t), \theta_G)]^T, \quad x(t) \in \mathcal{S}, \quad t \geq 0, \end{aligned}$$

and the feedback adversary strategy

$$(4.10) \quad \begin{aligned} & a(t) = a^*(x(t), \theta_a) \\ & \triangleq \frac{1}{2} R_a^{-1}(x(t)) [L_a(x(t)) + V'(x(t))K(x(t), \theta_K)]^T, \quad x(t) \in \mathcal{S}, \quad t \geq 0, \end{aligned}$$

the zero solution $x(t) \equiv 0$ to (3.2) is safely predefined-time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$. Furthermore, if $x_0 \in \mathcal{S}$, then the pair of feedback strategies $(u^*(\cdot), d^*(\cdot))$ is the Nash equilibrium of the two-player zero-sum game in the sense of (3.14) and the Nash value is given by (3.15).

Proof. The proof follows immediately from Theorem 3.2 with $\mathcal{D} = \mathbb{R}^n$, $U = \mathbb{R}^{m_u}$, $A = \mathbb{R}^{m_a}$, $\theta_s = [\theta_f^T, \theta_G^T, \theta_K^T]^T$, $F(x, \theta_s, u, a) = f(x, \theta_f) + G(x, \theta_G)u + K(x, \theta_K)a$, and $r(x, u, a) = L(x) + L_u(x)u + L_a(x)a + u^T R_u(x)u - a^T R_a(x)a$. Specifically, the

Hamiltonian function becomes

$$\begin{aligned}
& H(x, \theta_s, u, a, V'^T(x)) \\
&= L(x) + L_u(x)u + L_a(x)a + u^T R_u(x)u - a^T R_a(x)a \\
&\quad + V'(x) \left(f(x, \theta_f) + G(x, \theta_G)u + K(x, \theta_K)a \right), \\
(4.11) \quad & (x, \theta_s, u, a) \in \mathcal{S} \times \mathbb{R}^{N_s} \times \mathbb{R}^{m_u} \times \mathbb{R}^{m_a}.
\end{aligned}$$

Substituting $u = u^*(x, \theta_c)$ and $a = a^*(x, \theta_a)$, given by (4.9) and (4.10), into the Hamiltonian function (4.11) yields (4.4), and hence,

$$H(x, \theta_s, u^*(x, \theta_c), a^*(x, \theta_a), V'^T(x)) = 0, \quad x \in \mathcal{S},$$

which implies (3.5).

Next, using (3.5), the Hamiltonian function (4.11) can be written as

$$\begin{aligned}
H(x, \theta_s, u, a, V'^T(x)) &= H(x, \theta_s, u, a, V'^T(x)) \\
&\quad - H(x, \theta_s, u^*(x, \theta_c), a^*(x, \theta_a), V'^T(x)) \\
&= r(x, u, a) + V'(x) \left(f(x, \theta_f) + G(x, \theta_G)u \right. \\
&\quad \left. + K(x, \theta_K)a \right) - r(x, u^*(x, \theta_c), a^*(x, \theta_a)) \\
&\quad - V'(x) \left(f(x, \theta_f) + G(x, \theta_G)u^*(x, \theta_c) \right. \\
&\quad \left. + K(x, \theta_K)a^*(x, \theta_a) \right) \\
&= \left(L_a(x) + V'(x)K(x, \theta_K) \right) \left(a - a^*(x, \theta_a) \right) \\
&\quad - a^T R_a(x)a + a^{*T}(x, \theta_c) R_a(x) a^*(x, \theta_a) \\
&\quad + \left(L_u(x) + V'(x)G(x, \theta_G) \right) \left(u - u^*(x, \theta_c) \right) \\
&\quad + u^T R_u(x)u - u^{*T}(x, \theta_c) R_u(x) u^*(x, \theta_c),
\end{aligned}$$

which, using (4.9) and (4.10), yields

$$\begin{aligned}
H(x, \theta_s, u, a, V'^T(x)) &= 2a^{*T}(x, \theta_a) R_a(x) \left(a - a^*(x, \theta_a) \right) \\
&\quad - a^T R_a(x)a + a^{*T}(x, \theta_c) R_a(x) a^*(x, \theta_a), \\
&\quad - 2u^{*T}(x, \theta_c) R_u(x) \left(u - u^*(x, \theta_c) \right) \\
&\quad + u^T R_u(x)u - u^{*T}(x, \theta_c) R_u(x) u^*(x, \theta_c) \\
&= - \left(a - a^*(x, \theta_a) \right)^T R_a(x) \left(a - a^*(x, \theta_a) \right) \\
&\quad + \left(u - u^*(x, \theta_c) \right)^T R_u(x) \left(u - u^*(x, \theta_c) \right).
\end{aligned}$$

Now, since $R_a(x) > 0$, $x \in \mathcal{S}$, and $R_u(x) > 0$, $x \in \mathcal{S}$, it follows that

$$\begin{aligned}
H(x, \theta_s, u^*(x, \theta_c), a, V'^T(x)) &= - \left(a - a^*(x, \theta_a) \right)^T R_a(x) \left(a - a^*(x, \theta_a) \right) \\
&\leq 0, \quad (x, a) \in \mathcal{S} \times \mathbb{R}^{m_a},
\end{aligned}$$

and

$$\begin{aligned} H(x, \theta_s, u, a^*(x, \theta_a), V'^T(x)) &= \left(u - u^*(x, \theta_c)\right)^T R_u(x) \left(u - u^*(x, \theta_c)\right) \\ &\geq 0, \quad (x, u) \in \mathcal{S} \times \mathbb{R}^{m_u}, \end{aligned}$$

which imply (3.6) and (3.7).

Next, since $x = 0$ is a local minimum of the continuously differentiable $V(\cdot)$, we have $V'(0) = 0$. Hence, it follows from (4.9) and (4.10) that $u^*(0, \theta_c) = 0$ and $a^*(0, \theta_a) = 0$, which verify (3.8) and (3.9). Finally, (4.5)–(4.8) are equivalent to (3.10)–(3.13) with $u^*(x, \theta_c)$ and $a^*(x, \theta_a)$ given by (4.9) and (4.10). Now, the result is an immediate consequence of Theorem 3.2. $\blacksquare$

REMARK 6. *Since the Hamiltonian function (4.11) is convex in u and concave in a for every $x \in \mathcal{S}$, we derive the saddle point feedback strategies (4.9) and (4.10) by setting $\frac{\partial H}{\partial u} = 0$ and $\frac{\partial H}{\partial a} = 0$.* $\square$

The robust optimal safe predefined time stabilization problem is equivalent to solving the steady-state HJI equation (4.4), which is generally difficult to solve, except for special cases. To circumvent the complexity in solving the steady-state HJI equation, we consider a *robust inverse optimal feedback control problem* [15, 31, 30], wherein we parameterize a class of safely predefined time stabilizing pairs of feedback strategies that constitute a saddle point of a *derived* cost functional instead of seeking a *given* cost functional saddle point, thus providing flexibility in specifying the adversarially robust optimal controller.

COROLLARY 4.2. *Consider the parameter-dependent nonlinear affine dynamical system (4.1) with cost functional (3.3). Let the controller $u(\cdot)$ be a player that minimizes the cost functional (3.3) and let the adversary $a(\cdot)$ be a player that maximizes the cost functional (3.3). Let $\mathcal{S} \subset \mathbb{R}^n$ be a set of admissible states with $0 \in \mathcal{S}$ and let $T_p > 0$ be a predefined time. For every $x_0 \in \mathcal{S}$, let $\mathcal{F}(x_0, \mathcal{S}, T_p) \subset \mathcal{U} \times \mathcal{D}$ be the set of safely predefined time stabilizing pairs of feedback strategies and suppose that $\mathcal{F}(x_0, \mathcal{S}, T_p)$ is nonempty. Assume that there exist a continuously differentiable function $V : \mathcal{S} \rightarrow \mathbb{R}$, continuous functions $L_u : \mathcal{S} \rightarrow \mathbb{R}^{1 \times m_u}$ and $L_a : \mathcal{S} \rightarrow \mathbb{R}^{1 \times m_a}$, continuous positive-definite matrix functions $R_u : \mathcal{S} \rightarrow \mathbb{R}^{m_u \times m_u}$ and $R_a : \mathcal{S} \rightarrow \mathbb{R}^{m_a \times m_a}$, system parameter vectors $\theta_f \in \mathbb{R}^{N_f}$, $\theta_G \in \mathbb{R}^{N_G}$, and $\theta_K \in \mathbb{R}^{N_K}$, and real numbers $\alpha, \beta, p, q, r > 0$ such that $pr < 1$, $qr > 1$, and (4.5)–(4.8) hold. If either $\mathcal{S}$ is bounded or both $\mathcal{S}$ is unbounded and $V(\cdot)$ is coercive, then with the feedback control strategy (4.9) and the feedback adversary strategy (4.10), the zero solution $x(t) \equiv 0$ to (3.2) is safely predefined-time stable with predefined time T_p with respect to the set of admissible states $\mathcal{S}$, the cost functional (3.3), with*

$$(4.12) \quad L(x) = u^{*T}(x, \theta_c) R_u(x) u^*(x, \theta_c) - V'(x) f(x, \theta_f) - d^{*T}(x, \theta_a) R_a(x) a^*(x, \theta_a), \quad x \in \mathcal{S},$$

has a saddle point in the sense of (3.14), and the Nash value is given by (3.15).

Proof. The proof is similar to the proof of Corollary 4.1 and hence is omitted. $\blacksquare$

REMARK 7. *The following observations are important.*

1) *The function $L(x)$ in the running cost (4.2) given by (4.12) explicitly depends on the nonlinear system dynamics, the safely predefined time stabilizing Lyapunov function of the closed-loop system, and the feedback Nash equilibrium of the two-player*

zero-sum game. Note that the coupling is introduced via the steady-state HJI equation (4.4).

2) In light of (4.9) and (4.10), the vector functions $L_u(x)$ and $L_a(x)$ in the running cost (4.2) provide flexibility in the controller and adversary synthesis since $L_u(x)$ and $L_a(x)$ are arbitrary functions of $x \in \mathcal{S}$ subject to condition (4.8), a stability condition in Theorem 4.1. $\square$

In light of the above observations, note that by varying the parameters in the barrier Lyapunov function and the running cost, we can characterize a family of safely predefined-time stabilizing pairs of feedback strategies that can satisfy closed-loop system response requirements.

5. PINNs for Robust Optimal Safe Control. Unlike the robust *inverse* optimal safe stabilization problem, for which the Nash value is known, the *robust optimal safe predefined-time stabilization* problem is equivalent to solving the steady-state HJI equation (4.4) subject to the constraints (4.5)–(4.8), which is generally difficult to solve, except for special cases.

Building on the results of [36], we design a *physics-informed learning* framework to approximate the safely predefined-time stabilizing solution $V(\cdot)$ of the steady-state HJI equation (4.4). Specifically, invoking the universal approximation property of deep neural networks [19], we introduce a surrogate model $\hat{V}(\cdot, w)$ with $w \in \mathbb{R}^N$ representing the trainable model parameters. To ensure the satisfaction of the constraints (4.5)–(4.7), let

$$(5.1) \quad \hat{V}(x, w) = h(V_{\text{NN}}(x, w))B(x), \quad (x, w) \in \mathcal{S} \times \mathbb{R}^N,$$

where $h : \mathbb{R} \rightarrow (0, \infty)$ is a user-defined continuously differentiable function, $V_{\text{NN}} : \mathcal{S} \times \mathbb{R}^N \rightarrow \mathbb{R}$ is a standard fully-connected neural network, and $B : \mathcal{S} \rightarrow \mathbb{R}$ is a user-defined continuously differentiable function satisfying $B(0) = 0$, $B(x) > 0$, $x \in \mathcal{S} \setminus \{0\}$, and $B(x) \rightarrow \infty$ as $x \rightarrow \partial\mathcal{S}$. If $\mathcal{S}$ is unbounded, then $B(\cdot)$ is additionally coercive. The construction of $h(\cdot)$ and $B(\cdot)$ for each numerical example is provided in Section 6.

To learn an approximation $\hat{V}(\cdot, w)$ of $V(\cdot)$, we randomly sample M points in $\mathcal{S}$ and let $S_{\text{col}} \triangleq \{x_1, \dots, x_M\} \subset \mathcal{S}$ denote the set of *collocation points*. The learning procedure is formulated as a *constrained* optimization problem, where, at the collocation points, the loss function $\mathcal{E}(\cdot)$ penalizes violations of the steady-state HJI equation (4.4) and the constraint function $l(\cdot, \cdot)$ enforces the differential inequality (4.8) to guarantee safe predefined-time stability. Specifically, the constrained optimization problem is of the form

$$(5.2) \quad \begin{aligned} & \min_{w \in \mathbb{R}^N} \mathcal{E}(w) \\ & \text{subject to } l(x, w) \leq 0, \quad x \in S_{\text{col}}, \end{aligned}$$

where the loss function $\mathcal{E}(\cdot)$ and the constraint function $l(\cdot, \cdot)$ are defined by

$$(5.3) \quad \begin{aligned} \mathcal{E}(w) \triangleq & \sum_{x \in S_{\text{col}}} \left| L(x) + \hat{V}_x(x, w)f(x, \theta_f) \right. \\ & - \frac{1}{4} \left(\hat{V}_x(x, w)G(x, \theta_G) + L_u(x) \right) R_u^{-1}(x) \left(\hat{V}_x(x, w)G(x, \theta_G) + L_u(x) \right)^T \\ & \left. + \frac{1}{4} \left(\hat{V}_x(x, w)K(x, \theta_K) + L_a(x) \right) R_a^{-1}(x) \left(\hat{V}_x(x, w)K(x, \theta_K) + L_a(x) \right)^T \right|^2, \\ & w \in \mathbb{R}^N, \end{aligned}$$

and

$$\begin{aligned}
l(x, w) \triangleq & \hat{V}_x(x, w) \left(f(x, \theta_f) - \frac{1}{2} G(x, \theta_G) R_u^{-1}(x) L_u^T(x) \right. \\
& + \frac{1}{2} K(x, \theta_K) R_a^{-1}(x) L_a^T(x) \\
& - \frac{1}{2} G(x, \theta_G) R_u^{-1}(x) G^T(x, \theta_G) \hat{V}_x^T(x, w) \\
& + \frac{1}{2} K(x, \theta_K) R_a^{-1}(x) K^T(x, \theta_K) \hat{V}_x^T(x, w) \Big) \\
& + \frac{\gamma}{T_p} \left(\alpha \hat{V}^p(x, w) + \beta \hat{V}^q(x, w) \right)^r, \quad (x, w) \in \mathcal{S}_{\text{col}} \times \mathbb{R}^N.
\end{aligned} \tag{5.4}$$

The constrained optimization problem (5.2) can be numerically addressed using the *augmented Lagrangian method* (see, for example, [40] and [9]). In particular, the constrained optimization problem (5.2) is converted into a sequence of unconstrained optimization subproblems through an iterative procedure. In each iteration, the objective function of the unconstrained optimization subproblem is the sum of the loss function $\mathcal{E}(\cdot)$ and a penalty term for the constraints. Hence, the optimization subproblem at iteration k takes the form

$$(5.5) \quad \min_{w \in \mathbb{R}^N} \mathcal{E}_k(w),$$

where $\mathcal{E}_k(\cdot)$ is the loss function at iteration k given by

$$(5.6) \quad \mathcal{E}_k(w) \triangleq \mathcal{E}(w) + \sum_{x \in \mathcal{S}_{\text{col}}} \left(\mu_{k-1} \mathbb{1}_{\{l(x, w) \geq 0 \vee \lambda_{k-1}(x) > 0\}} l^2(x, w) + \lambda_{k-1}(x) l(x, w) \right),$$

with $\vee$ representing the or operator.

For every iteration $k \in \mathbb{Z}_+$, the Lagrangian multipliers μ_k and $\lambda_k(\cdot)$ in (5.6) are updated as

$$(5.7) \quad \mu_k = \delta \mu_{k-1},$$

$$(5.8) \quad \lambda_k(x) = \max\{0, \lambda_{k-1}(x) + 2\mu_{k-1} l(x, w)\}, \quad x \in \mathcal{S}_{\text{col}},$$

with $\delta, \mu_0 \in \mathbb{R}$ being tunable hyperparameters, and $\lambda_0(x)$ is initialized to be 0 for every $x \in \mathcal{S}_{\text{col}}$. The optimization subproblem (5.5) can be solved using modern gradient-based or (quasi-) Newton-based numerical solvers, such as Adam [25] or L-BFGS [53].

Let w_k denote the minimizer of the optimization subproblem (5.5) at iteration k . Substituting w_k into (5.1) yields $\hat{V}(\cdot, w_k)$, an approximation of the value function. Furthermore, substituting $\hat{V}(\cdot, w_k)$ into (4.9) and (4.10) yields the approximate Nash feedback control and adversary strategies $\hat{u}(\cdot, w_k)$ and $\hat{a}(\cdot, w_k)$, respectively. Algorithm 5.1 presents the pseudocode for the proposed adversarial physics-informed learning framework for robust optimal safe control, executed for a given number of iterations $\mathcal{K}$. The architecture of Algorithm 5.1 is shown in Fig. 2.

REMARK 8. Note that the steady-state HJI equation (4.4) may have multiple solutions. However, the Nash value $V(\cdot)$ is the unique safely predefined-time stabilizing solution. Hence, when using PINNs to solve the robust optimal safe predefined-time stabilization problem, enforcing the stability conditions (4.5)-(4.8) is imperative. Conditions (4.5)-(4.7) are satisfied by the design of the PINN architecture, while (4.8) is imposed as a constraint in the learning procedure. $\square$

Algorithm 5.1 Training procedure of Adversarial PINN

Hyperparameters: $\alpha, \beta, \gamma, T_p, p, q, r, \delta, \mu_0$

Input: Collocation points $S_{\text{col}} \triangleq \{x_1, \dots, x_M\} \subset \mathcal{S}$

Output: Nash value $\hat{V}(\cdot, w_K)$ and Nash equilibrium $(\hat{u}(\cdot, w_K), \hat{a}(\cdot, w_K))$

```

1: procedure
2:    $\lambda_0(x) \leftarrow 0$  for every  $x \in S_{\text{col}}$ 
3:   Initialize network parameters  $w_0 \in \mathbb{R}^N$ 
4:   for  $k = 1 \dots K$  do
5:      $\tilde{\mathcal{E}} \leftarrow \mathcal{E}(w_{k-1})$   $\triangleright (5.3)$ 
6:      $\tilde{l}(x) \leftarrow l(x, w_{k-1})$   $\triangleright (5.4)$ 
7:      $\mathcal{E}_k \leftarrow \tilde{\mathcal{E}} + \sum_{x \in S_{\text{col}}} \left( \mu_{k-1} \mathbb{I}_{\{\tilde{l}(x) \geq 0 \vee \lambda_{k-1}(x) > 0\}} \tilde{l}^2(x) + \lambda_{k-1}(x) \tilde{l}(x) \right)$ 
8:      $w_k \leftarrow \arg \min_w \mathcal{E}_k$ 
9:      $\mu_k \leftarrow \delta \mu_{k-1}$ 
10:     $\lambda_k(x) \leftarrow \max\{0, \lambda_{k-1}(x) + 2\mu_{k-1}\tilde{l}(x)\}$ 
        for every  $x \in S_{\text{col}}$ 
11:   end for
12: end procedure
  
```

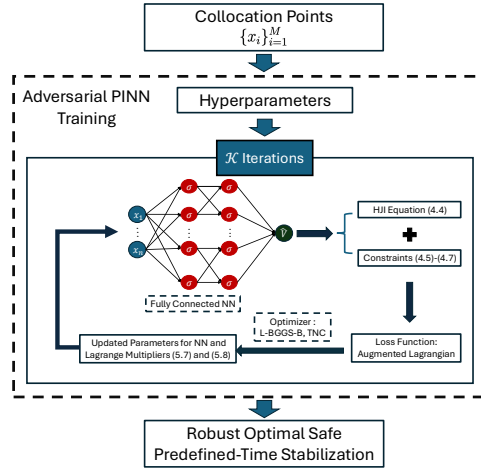


FIG. 2. Adversarial physics-informed learning architecture to solve the robust optimal safe predefined-time stabilization problem. A set of collocation points $\{x_i\}_{i=1}^M$ is randomly sampled in $\mathcal{S}$. The training is formulated as a constrained optimization problem: at each collocation point, the loss function penalizes deviations from the HJI equation (4.4), while the constraint function enforces the differential inequality (4.8) to ensure safe predefined-time stability. Constraints (4.5)-(4.7) are satisfied by the design of the PINN architecture. This constrained optimization problem is solved iteratively using the augmented Lagrangian method, with the Lagrange multipliers updated at each iteration according to (5.7) and (5.8).

6. Illustrative Numerical Examples. In this section, we present two numerical examples to illustrate the developed robust optimal safe predefined time control and physics-informed learning framework.

Consider the nonlinear affine dynamical system given by

$$(6.1) \quad \dot{x}_1(t) = -\min\{0, x_1(t)\} + u_1(t) + a_1(t), \quad x_1(0) = x_{10}, \quad t \geq 0,$$

$$(6.2) \quad \dot{x}_2(t) = -\max\{0, x_2(t)\} + u_2(t) + a_2(t), \quad x_2(0) = x_{20}.$$

Note that (6.1) and (6.2) can be cast in the form of (4.1) with $n = 2$, $m_u = 2$, $m_a = 2$, $x = [x_1, x_2]^T$, $u = [u_1, u_2]^T$, $a = [a_1, a_2]^T$, $f(x, \theta_f) = -[\min\{0, x_1\}, \max\{0, x_2\}]^T$, $G(x, \theta_G) = I_2$, and $K(x, \theta_K) = I_2$.

Let $s : \mathbb{R}^2 \rightarrow \mathbb{R}$ be a continuous function such that $s(0) > 0$. Define the set of admissible states $\mathcal{S}$ as a zero-strict superlevel set of $s(\cdot)$ given by

$$\mathcal{S} = \{x \in \mathbb{R}^2 : s(x) > 0\}.$$

Note that $\mathcal{S}$ is an open set with $0 \in \mathcal{S}$.

Next, given the set of admissible states $\mathcal{S}$ and a predefined time T_p , we use Corollary 4.2 to synthesize an inverse safely predefined time stabilizing Nash equilibrium $(u^*(\cdot), a^*(\cdot))$. In particular, we address the case of both a bounded and an unbounded set of admissible states.

6.1. Bounded Safe Set. Let $s(x) = \min\{s_1(x), s_2(x)\}$, $x \in \mathbb{R}^2$, where $s_1(x) \triangleq 1 - x_1^2$, $x \in \mathbb{R}^2$, and $s_2(x) \triangleq 1 - x_2^2$, $x \in \mathbb{R}^2$. Furthermore, let $V(x) = \frac{x_1^2}{2s_1(x)} + \frac{x_2^2}{2s_2(x)}$, $x \in \mathcal{S}$, be the Nash value. The terms of the running cost (4.2) are given by

$$\begin{aligned} L(x) &= \frac{1}{2} \left(\frac{[x_1]^{\gamma_1}}{s_1^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_1]^{\gamma_2}}{s_1^{\frac{\gamma_2-1}{2}}(x)} \right)^2 + \frac{1}{2} \left(\frac{[x_2]^{\gamma_1}}{s_2^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_2]^{\gamma_2}}{s_2^{\frac{\gamma_2-1}{2}}(x)} \right)^2 \\ &\quad + \frac{x_1}{s_1(x)} \left(1 + \frac{x_1^2}{s_1(x)} \right) \min\{0, x_1\} + \frac{x_2}{s_2(x)} \left(1 + \frac{x_2^2}{s_2(x)} \right) \max\{0, x_2\}, \quad x \in \mathcal{S}, \\ L_u(x) = L_a(x) &= \begin{bmatrix} \frac{[x_1]^{\gamma_1}}{s_1^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_1]^{\gamma_2}}{s_1^{\frac{\gamma_2-1}{2}}(x)} \\ \frac{[x_2]^{\gamma_1}}{s_2^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_2]^{\gamma_2}}{s_2^{\frac{\gamma_2-1}{2}}(x)} \end{bmatrix}^T \\ &\quad - \left[\frac{x_1}{s_1(x)} \left(1 + \frac{x_1^2}{s_1(x)} \right), \frac{x_2}{s_2(x)} \left(1 + \frac{x_2^2}{s_2(x)} \right) \right], \quad x \in \mathcal{S}, \end{aligned}$$

$R_u(x) = \frac{1}{4}I_2$, $x \in \mathcal{S}$, and $R_a(x) = \frac{1}{2}I_2$, $x \in \mathcal{S}$, where $\gamma_1 \in (0, 1)$ and $\gamma_2 > 1$. Hence, with $\theta_c = [\gamma_1, \gamma_2]^T$ and $\theta_a = [\gamma_1, \gamma_2]^T$, the inverse Nash equilibrium is given by

$$(6.3) \quad u^*(x, \theta_c) = -2 \begin{bmatrix} \frac{[x_1]^{\gamma_1}}{s_1^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_1]^{\gamma_2}}{s_1^{\frac{\gamma_2-1}{2}}(x)} \\ \frac{[x_2]^{\gamma_1}}{s_2^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_2]^{\gamma_2}}{s_2^{\frac{\gamma_2-1}{2}}(x)} \end{bmatrix}, \quad x \in \mathcal{S},$$

and

$$(6.4) \quad a^*(x, \theta_a) = \begin{bmatrix} \frac{[x_1]^{\gamma_1}}{s_1^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_1]^{\gamma_2}}{s_1^{\frac{\gamma_2-1}{2}}(x)} \\ \frac{[x_2]^{\gamma_1}}{s_2^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x_2]^{\gamma_2}}{s_2^{\frac{\gamma_2-1}{2}}(x)} \end{bmatrix}, \quad x \in \mathcal{S}.$$

Next, computing the time derivative of $V(x)$ along the trajectories of (6.1) and (6.2) with $u = u^*(x, \theta_c)$ and $a = a^*(x, \theta_a)$ given by (6.3) and (6.4), we obtain

$$\begin{aligned} \dot{V}(x) &= - \left(\frac{|x_1|^{\gamma_1+1}}{s_1^{\frac{\gamma_1+1}{2}}(x)} + \frac{|x_1|^{\gamma_2+1}}{s_1^{\frac{\gamma_2+1}{2}}(x)} \right) - \left(\frac{|x_2|^{\gamma_1+1}}{s_2^{\frac{\gamma_1+1}{2}}(x)} + \frac{|x_2|^{\gamma_2+1}}{s_2^{\frac{\gamma_2+1}{2}}(x)} \right) \\ &\quad - \frac{x_1}{s_1(x)} \min\{0, x_1\} - \frac{x_2}{s_2(x)} \max\{0, x_2\} \\ &\quad - \frac{x_1^2}{s_1^2(x)} \left(\min\{0, x_1\} x_1 + \frac{|x_1|^{\gamma_1+1}}{s_1^{\frac{\gamma_1-1}{2}}(x)} + \frac{|x_1|^{\gamma_2+1}}{s_1^{\frac{\gamma_2-1}{2}}(x)} \right) \\ &\quad - \frac{x_2^2}{s_2^2(x)} \left(\max\{0, x_2\} x_2 + \frac{|x_2|^{\gamma_1+1}}{s_2^{\frac{\gamma_1-1}{2}}(x)} + \frac{|x_2|^{\gamma_2+1}}{s_2^{\frac{\gamma_2-1}{2}}(x)} \right) \\ &\leq - \left(\frac{|x_1|^{\gamma_1+1}}{s_1^{\frac{\gamma_1+1}{2}}(x)} + \frac{|x_2|^{\gamma_1+1}}{s_2^{\frac{\gamma_1+1}{2}}(x)} \right) - \left(\frac{|x_1|^{\gamma_2+1}}{s_1^{\frac{\gamma_2+1}{2}}(x)} + \frac{|x_2|^{\gamma_2+1}}{s_2^{\frac{\gamma_2+1}{2}}(x)} \right), \quad x \in \mathcal{S}. \end{aligned}$$

Now, using $|y + \bar{y}|^{\delta_1} \leq |y|^{\delta_1} + |\bar{y}|^{\delta_1}$ for every $y, \bar{y} \in \mathbb{R}$ and $\delta_1 \in (0, 1)$, and $|y + \bar{y}|^{\delta_2} \leq 2^{\delta_2-1} (|y|^{\delta_2} + |\bar{y}|^{\delta_2})$ for every $y, \bar{y} \in \mathbb{R}$ and $\delta_2 > 1$ [39], and letting $y = \frac{x_1^2}{2s_1(x)}$, $x \in \mathcal{S}$, $\delta_1 = \frac{\gamma_1+1}{2}$, $\bar{y} = \frac{x_2^2}{2s_2(x)}$, $x \in \mathcal{S}$, and $\delta_2 = \frac{\gamma_2+1}{2}$, it follows that

$$\dot{V}(x) \leq -2^{\frac{\gamma_1+1}{2}} V^{\frac{\gamma_1+1}{2}}(x) - 2V^{\frac{\gamma_2+1}{2}}(x), \quad x \in \mathcal{S},$$

which implies that (4.8) holds with $\gamma = T_p$, $\alpha = 2^{\frac{\gamma_1+1}{2}}$, $\beta = 2$, $p = \frac{\gamma_1+1}{2}$, $q = \frac{\gamma_2+1}{2}$, and $r = 1$. Hence, by Corollary 4.2, the zero solution $x(t) \equiv 0$ of the closed-loop system is safely predefined-time stable.

Let $T_p = 3.4295$ so that $\gamma_1 = 0.5$ and $\gamma_2 = 1.5$. We use Algorithm 5.1 to learn the Nash value $V(\cdot)$, the Nash feedback control strategy $u^*(\cdot, \theta_c)$, and the Nash feedback adversary strategy $a^*(\cdot, \theta_a)$. Specifically, for the PINN surrogate $\hat{V}(x, w)$ defined by (5.1), we set $h(x) = e^x$, $x \in \mathbb{R}$, and $B(x) = \frac{x_1^2}{1-|x_1|} + \frac{x_2^2}{1-|x_2|}$, $x \in \mathcal{S}$. The neural network $V_{\text{NN}}(\cdot, \cdot)$ is implemented as a fully connected neural network with six hidden layers, each comprising 100 neurons, and employs the hyperbolic tangent activation function $\tanh(\cdot)$. The number of collocation points is set to $M = 10^4$. The hyperparameters associated with the augmented Lagrangian method are selected as $\mu_0 = 10^{-4}$ and $\delta = 2$. At each iteration of Algorithm 5.1, the optimization subproblem (5.5) is solved using the L-BFGS optimizer.

Fig. 3 shows the approximate Nash value $\hat{V}$, the approximate Nash control strategy $\hat{u}$, and the approximate Nash adversary strategy $\hat{a}$ obtained from Algorithm 5.1, along with the exact Nash value V , the exact Nash control strategy u^* , and the exact

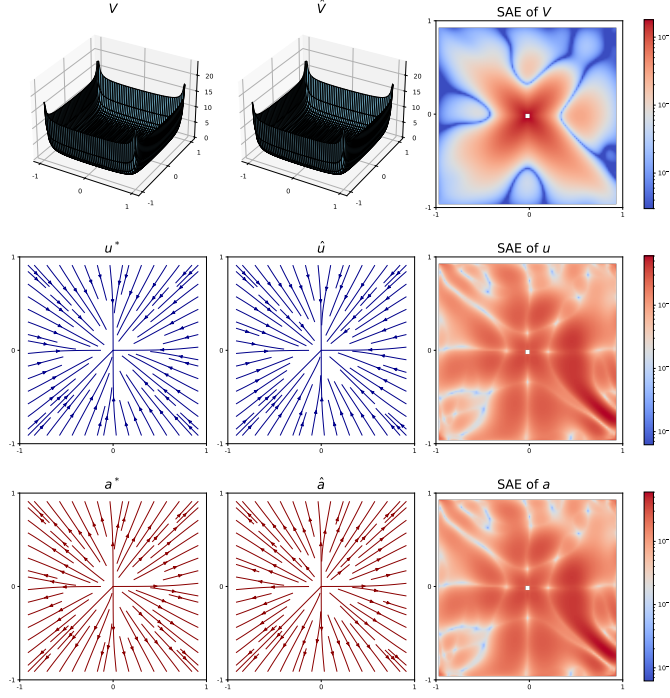


FIG. 3. Nash value, Nash feedback control strategy, and Nash feedback adversary strategy for the bounded safe set. **(Left)** Exact Nash value V , exact Nash control strategy u^* , and exact Nash adversary strategy a^* . **(Middle)** Approximate Nash value $\hat{V}$, approximate Nash control strategy $\hat{u}$, and approximate Nash adversary strategy $\hat{a}$. **(Right)** Symmetric absolute error for learning the Nash value, the Nash control strategy, and the Nash adversary strategy.

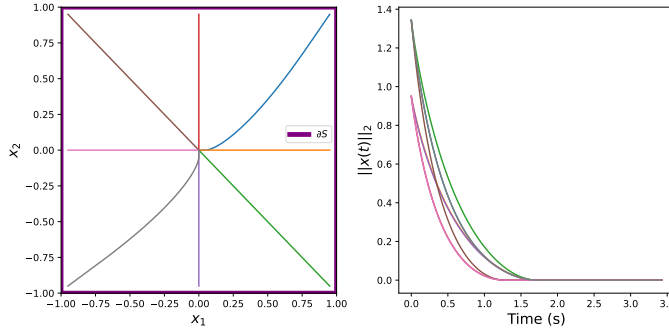


FIG. 4. Robust optimal safe predefined-time stabilization for the bounded safe set. **(Left)** Closed-loop state trajectories $x(t)$, $t \geq 0$, starting from different initial conditions near the boundary ∂S of the safe set. **(Right)** Time evolution of the Euclidean norm $\|x(t)\|_2$, $t \geq 0$, of each trajectory shown in the left plot. In both plots, trajectories starting from the same initial condition are marked with the same color.

Nash adversary strategy a^* . To evaluate the approximation accuracy of our algorithm, we define, for every $x \in S_{\text{col}}$, the symmetric absolute error (SAE) for the approximate Nash value $\hat{V}$, the approximate Nash control strategy $\hat{u}$, and the approximate Nash

adversary strategy $\hat{a}$ as

$$\begin{aligned}\text{SAE}(\hat{V}(x, w_{\mathcal{K}}), V(x)) &\triangleq \frac{|\hat{V}(x, w_{\mathcal{K}}) - V(x)|}{|\hat{V}(x, w_{\mathcal{K}})| + |V(x)|}, \\ \text{SAE}(\hat{u}(x, w_{\mathcal{K}}), u^*(x, \theta_c)) &\triangleq \frac{\|\hat{u}(x, w_{\mathcal{K}}) - u^*(x, \theta_c)\|_1}{\|\hat{u}(x, w_{\mathcal{K}})\|_1 + \|u^*(x, \theta_c)\|_1},\end{aligned}$$

and

$$\text{SAE}(\hat{a}(x, w_{\mathcal{K}}), a^*(x, \theta_a)) \triangleq \frac{\|\hat{a}(x, w_{\mathcal{K}}) - a^*(x, \theta_a)\|_1}{\|\hat{a}(x, w_{\mathcal{K}})\|_1 + \|a^*(x, \theta_a)\|_1}.$$

As shown in Fig. 3, the SAE between the exact and the approximate solution to the game learned by Algorithm 5.1 is low, verifying the algorithm's efficacy. The left plot of Fig. 4 shows the closed-loop state trajectories starting from different initial conditions near the boundary $\partial\mathcal{S}$ of the safe set $\mathcal{S}$. Note that all trajectories remain within $\mathcal{S}$ and converge to the origin. The right plot of Fig. 4 shows the time evolution of the Euclidean norm of each trajectory. Note the predefined time convergence to the origin, which confirms the settling-time function bound $T(x_0, \theta_c, \theta_a) \leq 3.4259$, $x_0 \in \mathcal{S}$.

6.2. Unbounded Safe Set. Let $s(x) = 1 - x_2^2$, $x \in \mathbb{R}^2$, and let $V(x) = \frac{\|x\|_2^2}{2s(x)}$, $x \in \mathcal{S}$, be the Nash value, whereas the terms composing the running cost (4.2) are chosen as

$$\begin{aligned}L(x) &= \frac{1}{2} \left\| \frac{[x]^{\gamma_1}}{s^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x]^{\gamma_2}}{s^{\frac{\gamma_2-1}{2}}(x)} \right\|_2^2 \\ &\quad + \frac{x_1}{s(x)} \min\{0, x_1\} + \frac{x_2}{s(x)} \left(1 + \frac{\|x\|_2^2}{s(x)} \right) \max\{0, x_2\}, \quad x \in \mathcal{S}, \\ L_u(x) &= L_a(x) \\ &= \left(\frac{[x]^{\gamma_1}}{s^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x]^{\gamma_2}}{s^{\frac{\gamma_2-1}{2}}(x)} \right)^T - \left[\frac{x_1}{s(x)}, \frac{x_2}{s(x)} \left(1 + \frac{\|x\|_2^2}{s(x)} \right) \right], \quad x \in \mathcal{S},\end{aligned}$$

$R_u(x) = \frac{1}{4}I_2$, $x \in \mathcal{S}$, and $R_a(x) = \frac{1}{2}I_2$, $x \in \mathcal{S}$, where $\gamma_1 \in (0, 1)$ and $\gamma_2 > 1$. Hence, with $\theta_c = [\gamma_1, \gamma_2]^T$ and $\theta_a = [\gamma_1, \gamma_2]^T$, the inverse Nash equilibrium is given by

$$(6.5) \quad u^*(x, \theta_c) = -2 \left(\frac{[x]^{\gamma_1}}{s^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x]^{\gamma_2}}{s^{\frac{\gamma_2-1}{2}}(x)} \right), \quad x \in \mathcal{S},$$

and

$$(6.6) \quad a^*(x, \theta_a) = \left(\frac{[x]^{\gamma_1}}{s^{\frac{\gamma_1-1}{2}}(x)} + \frac{[x]^{\gamma_2}}{s^{\frac{\gamma_2-1}{2}}(x)} \right), \quad x \in \mathcal{S}.$$

Next, evaluating the time derivative of $V(x)$ along the trajectories of (6.1) and (6.2) with $u = u^*(x, \theta_c)$ and $a = a^*(x, \theta_a)$ given by (6.5) and (6.6) yields

$$\begin{aligned}
\dot{V}(x) &= - \left(\frac{\|x\|_{\gamma_1+1}^{\gamma_1+1}}{s^{\frac{\gamma_1+1}{2}}(x)} + \frac{\|x\|_{\gamma_2+1}^{\gamma_2+1}}{s^{\frac{\gamma_2+1}{2}}(x)} \right) - \frac{1}{s(x)} (x_1 \min\{0, x_1\} + x_2 \max\{0, x_2\}) \\
&\quad - \frac{\|x\|_2^2}{s^2(x)} \left(\max\{0, x_2\} x_2 + \frac{|x_2|^{\gamma_1+1}}{s^{\frac{\gamma_1-1}{2}}(x)} + \frac{|x_2|^{\gamma_2+1}}{s^{\frac{\gamma_2-1}{2}}(x)} \right) \\
(6.7) \quad &\leq - \left(\frac{\|x\|_{\gamma_1+1}^{\gamma_1+1}}{s^{\frac{\gamma_1+1}{2}}(x)} + \frac{\|x\|_{\gamma_2+1}^{\gamma_2+1}}{s^{\frac{\gamma_2+1}{2}}(x)} \right), \quad x \in \mathcal{S}.
\end{aligned}$$

Now, by the monotonicity property of Hölder norms [7], we have $\|x\|_2 \leq \|x\|_{\gamma_1+1}$ since $2 > \gamma_1 + 1 > 0$. Moreover, by the equivalence of vector norms on $\mathbb{R}^n$ [7], it follows that $\|x\|_2 \leq 2^{\frac{\gamma_2-1}{2(\gamma_2+1)}} \|x\|_{\gamma_2+1}$ since $2 < \gamma_2 + 1$, and hence, (6.7) can be further bounded as

$$\dot{V}(x) \leq -2^{\frac{\gamma_1+1}{2}} V^{\frac{\gamma_1+1}{2}}(x) - 2V^{\frac{\gamma_2+1}{2}}(x), \quad x \in \mathcal{S},$$

which implies that (4.8) is satisfied with $\gamma = T_p$, $\alpha = 2^{\frac{\gamma_1+1}{2}}$, $\beta = 2$, $p = \frac{\gamma_1+1}{2}$, $q = \frac{\gamma_2+1}{2}$, and $r = 1$, and hence, by Corollary 4.2, the zero solution $x(t) \equiv 0$ of the closed-loop system is safely predefined-time stable.

Let $T_p = 3.4259$ so that $\gamma_1 = 0.5$ and $\gamma_2 = 1.5$. For our PINN architecture, we let $h(x) = \frac{1}{1+e^{-x}}$, $x \in \mathbb{R}$, which is the sigmoid function, and $B(x) = \frac{\|x\|_2^2}{\sqrt{2}-\sqrt{x_2^2+1}}$, $x \in \mathcal{S}$.

The neural network $V_{\text{NN}}(\cdot, \cdot)$ is implemented as a fully connected neural network with six hidden layers, each comprising 50 neurons, using the sigmoid activation function. The number of collocation points is set to $M = 10^4$. For the augmented Lagrangian method, the hyperparameters are set to $\mu_0 = 10^{-3}$ and $\delta = 1.5$. At every iteration of Algorithm 5.1, the optimization subproblem (5.5) is solved using the truncated Newton conjugate gradient (TNCG) optimizer (see, for example, [24] and [11]).

Fig. 5 shows the approximate Nash value $\hat{V}$, the approximate Nash control strategy $\hat{u}$, and the approximate Nash adversary strategy $\hat{a}$, along with the exact Nash value V , the exact Nash control strategy u^* , and the exact Nash adversary strategy a^* . Note that Algorithm 5.1 exhibits a low SAE. The left plot of Fig. 6 shows the closed-loop state trajectories starting from different initial conditions within the safe set $\mathcal{S}$. Note that all trajectories remain within $\mathcal{S}$ and converge to the origin. The right plot of Fig. 6 shows the time evolution of the Euclidean norm of each trajectory. Note the predefined time convergence to the origin, which verifies the settling-time function bound $T(x_0, \theta_c, \theta_a) \leq 3.4259$, $x_0 \in \mathcal{S}$.

7. Conclusion. In this paper, an adversarially robust optimal safe predefined-time stabilization problem is stated and shown to correspond to a two-player zero-sum differential game. Sufficient conditions are derived to characterize a safely predefined-time stabilizing saddle-point solution to the differential game. Specifically, safe predefined-time stability of the closed-loop system and Nash equilibrium are established via a barrier Lyapunov function that satisfies a certain differential inequality and the steady-state HJI equation. Given the intractability of the latter, an adversarial physics-informed learning algorithm is developed to learn the safely predefined-time stabilizing solution to the steady-state HJI equation. Future research will explore discrete-time extensions of the proposed framework.

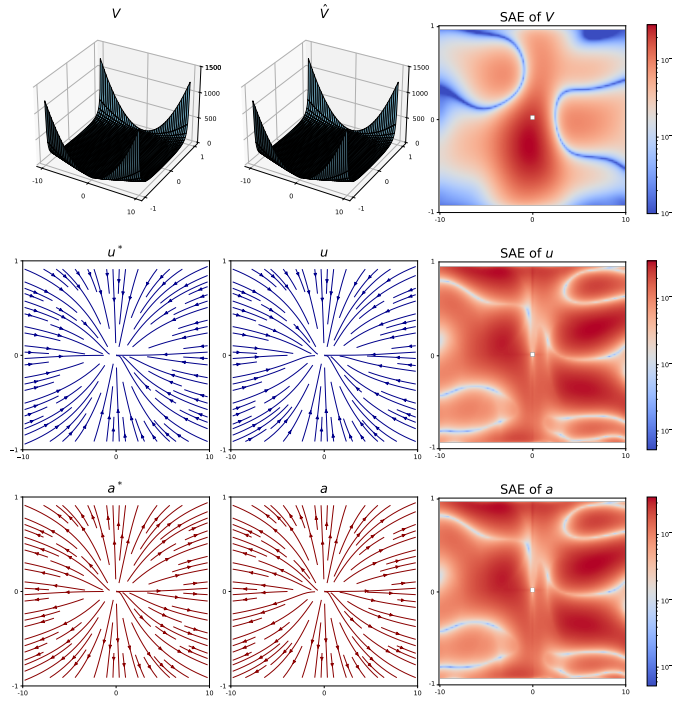


FIG. 5. Nash value, Nash feedback control strategy, and Nash feedback adversary strategy for the unbounded safe set. **(Left)** Exact Nash value V , exact Nash control strategy u^* , and exact Nash adversary strategy a^* . **(Middle)** Approximate Nash value $\hat{V}$, approximate Nash control strategy $\hat{u}$, and approximate Nash adversary strategy $\hat{a}$. **(Right)** Symmetric absolute error for learning the Nash value, the Nash control strategy, and the Nash adversary strategy.

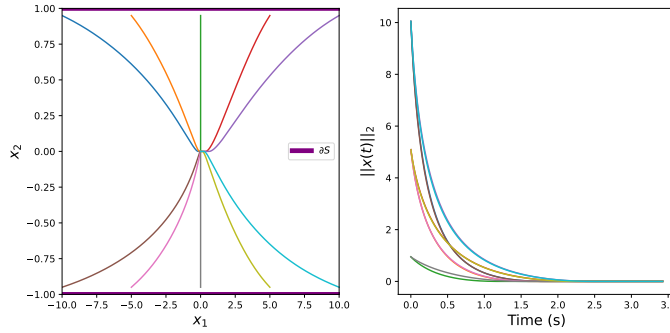


FIG. 6. Robust optimal safe predefined-time stabilization for the unbounded safe set. **(Left)** Closed-loop state trajectories $x(t)$, $t \geq 0$, starting from different initial conditions in the safe set $\mathcal{S}$. **(Right)** Time evolution of the Euclidean norm $\|x(t)\|_2$, $t \geq 0$, of each trajectory shown in the left plot. In both plots, trajectories starting from the same initial condition are marked with the same color.

REFERENCES

- [1] A. D. AMES, S. COOGAN, M. EGERSTEDT, G. NOTOMISTA, K. SREENATH, AND P. TABUADA, Control barrier functions: Theory and applications, in European Control Conference, 2019,

- pp. 3420–3431.
- [2] A. D. AMES, X. XU, J. W. GRIZZLE, AND P. TABUADA, *Control barrier function based quadratic programs for safety critical systems*, IEEE Transactions on Automatic Control, 62 (2016), pp. 3861–3876.
 - [3] J. A. BALL AND J. HELTON, *H/sup infinity/control for nonlinear plants: connections with differential games*, in Proceedings of the 28th IEEE Conference on Decision and Control,, IEEE, 1989, pp. 956–962.
 - [4] T. BAŞAR AND P. BERNHARD, *H-infinity optimal control and related minimax design problems: a dynamic game approach*, Springer Science & Business Media, New York, 2008.
 - [5] T. BAŞAR AND G. J. OLSDER, *Dynamic noncooperative game theory*, SIAM, Philadelphia, 1998.
 - [6] D. S. BERNSTEIN, *Nonquadratic cost and nonlinear feedback control*, International Journal of Robust and Nonlinear Control, 3 (1993), pp. 211–229.
 - [7] D. S. BERNSTEIN, *Scalar, Vector, and Matrix Mathematics: Theory, Facts, and Formulas*, Princeton University Press, Princeton, NJ, 2018.
 - [8] D. BERTSEKAS, *Nonlinear Programming*, Athena Scientific, Nashua, NH, 2011.
 - [9] D. P. BERTSEKAS, *Constrained optimization and Lagrange multiplier methods*, Academic press, Boston, 2014.
 - [10] S. P. BHAT AND D. S. BERNSTEIN, *Finite-time stability of continuous autonomous systems*, SIAM Journal on Control and Optimization, 38 (2000), pp. 751–766.
 - [11] J.-F. BONNANS, J. C. GILBERT, C. LEMARÉCHAL, AND C. A. SAGASTIZÁBAL, *Numerical optimization: theoretical and practical aspects*, Springer Science & Business Media, New York, 2006.
 - [12] P. BRAUN AND C. M. KELLETT, *Comment on “stabilization with guaranteed safety using control lyapunov–barrier function”*, Automatica, 122 (2020), p. 109225.
 - [13] M. COHEN AND C. BELTA, *Adaptive and Learning-Based Control of Safety-Critical Systems*, Springer Nature, Cham, 2023.
 - [14] F. FOTIADIS AND K. G. VAMVOUDAKIS, *A physics-informed neural networks framework to solve the infinite-horizon optimal control problem*, in 2023 62nd IEEE Conference on Decision and Control (CDC), IEEE, 2023, pp. 6014–6019.
 - [15] R. A. FREEMAN AND P. V. KOKOTOVIC, *Inverse optimality in robust stabilization*, SIAM Journal on Control and Optimization, 34 (1996), pp. 1365–1391.
 - [16] R. FURFARO, A. D’AMBROSIO, E. SCHIASSI, AND A. SCORSOGLIO, *Physics-informed neural networks for closed-loop guidance and control in aerospace systems*, in AIAA SCITECH 2022 Forum, 2022, p. 0361.
 - [17] W. M. HADDAD AND V. CHELLABOINA, *Nonlinear Dynamical Systems and Control: A Lyapunov-Based Approach*, Princeton University Press, Princeton, NJ, 2011.
 - [18] W. M. HADDAD AND A. L’AFFLITTO, *Finite-time stabilization and optimal feedback control*, IEEE Transactions on Automatic Control, 61 (2015), pp. 1069–1074.
 - [19] K. HORNIK, M. STINCHCOMBE, AND H. WHITE, *Universal approximation of an unknown mapping and its derivatives using multilayer feedforward networks*, Neural Networks, 3 (1990), pp. 551–560.
 - [20] D. JACOBSON, *On values and strategies for infinite-time linear quadratic games*, IEEE Transactions on Automatic Control, 22 (1977), pp. 490–491.
 - [21] M. JANKOVIC, *Robust control barrier functions for constrained stabilization of nonlinear systems*, Automatica, 96 (2018), pp. 359–367.
 - [22] E. JIMÉNEZ-RODRÍGUEZ, A. J. MUÑOZ-VÁZQUEZ, J. D. SÁNCHEZ-TORRES, M. DEFOORT, AND A. G. LOUKIANOV, *A lyapunov-like characterization of predefined-time stability*, IEEE Transactions on Automatic Control, 65 (2020), pp. 4922–4927.
 - [23] E. JIMÉNEZ-RODRÍGUEZ, J. D. SÁNCHEZ-TORRES, AND A. G. LOUKIANOV, *On optimal predefined-time stabilization*, International Journal of Robust and Nonlinear Control, 27 (2017), pp. 3620–3642.
 - [24] C. T. KELLEY, *Iterative methods for optimization*, SIAM, Philadelphia, 1999.
 - [25] D. P. KINGMA AND J. BA, *Adam: A method for stochastic optimization*, arXiv preprint arXiv:1412.6980, (2014).
 - [26] N.-M. T. KOKOLAKIS AND K. G. VAMVOUDAKIS, *Safe predefined-time stability and optimal feedback control: A lyapunov-based approach*, in 2024 American Control Conference (ACC), IEEE, 2024, pp. 3680–3685.
 - [27] N.-M. T. KOKOLAKIS, K. G. VAMVOUDAKIS, AND W. M. HADDAD, *Fixed-time learning for optimal feedback control*, in International Design Engineering Technical Conferences and Computers and Information in Engineering Conference, vol. 87387, American Society of Mechanical Engineers, 2023, p. V010T10A027.
 - [28] N.-M. T. KOKOLAKIS, Z. ZHANG, S. LIU, K. G. VAMVOUDAKIS, J. DARBON, AND G. E. KARNI-

- ADAKIS, *Safe physics-informed machine learning for optimal predefined-time stabilization: A lyapunov-based approach*, IEEE Transactions on Neural Networks and Learning Systems, (2025).
- [29] A. L'AFFLITTO, *Differential games, asymptotic stabilization, and robust optimal control of nonlinear systems*, in 2016 IEEE 55th Conference on Decision and Control (CDC), IEEE, 2016, pp. 1933–1938.
 - [30] A. L'AFFLITTO, *Differential games, continuous lyapunov functions, and stabilisation of nonlinear dynamical systems*, IET Control Theory & Applications, 11 (2017), pp. 2486–2496.
 - [31] A. L'AFFLITTO, *Differential games, finite-time partial-state stabilisation of nonlinear dynamical systems, and optimal robust control*, International Journal of Control, 90 (2017), pp. 1861–1878.
 - [32] A. L'AFFLITTO, *Differential games, partial-state stabilization, and model reference adaptive control*, J. Frankl. Inst., 354 (2017), pp. 456–478.
 - [33] F. L. LEWIS, D. VRABIE, AND V. L. SYRMOS, *Optimal Control*, John Wiley & Sons, Hoboken, NJ, 2012.
 - [34] D. LIBERZON, *Calculus of Variations and Optimal Control Theory: A Concise Introduction*, Princeton University Press, Princeton, NJ, 2011.
 - [35] D. J. LIMEBEER, B. D. ANDERSON, P. P. KHARGONEKAR, AND M. GREEN, *A game theoretic approach to h_∞ control for time-varying systems*, SIAM Journal on Control and Optimization, 30 (1992), pp. 262–283.
 - [36] L. LU, R. PESTOURIE, W. YAO, Z. WANG, F. VERDUGO, AND S. G. JOHNSON, *Physics-informed neural networks with hard constraints for inverse design*, SIAM Journal on Scientific Computing, 43 (2021), pp. B1105–B1132.
 - [37] E. MAGEIROU, *Values and strategies for infinite time linear quadratic games*, IEEE Transactions on Automatic Control, 21 (1976), pp. 547–550.
 - [38] P. MESTRES AND J. CORTÉS, *Optimization-based safe stabilizing feedback with guaranteed region of attraction*, IEEE Control Systems Letters, 7 (2022), pp. 367–372.
 - [39] D. S. MITRINOVIC AND P. M. VASIC, *Analytic Inequalities*, vol. 1, Springer, Berlin, Heidelberg, 1970.
 - [40] J. NOCEDAL AND S. J. WRIGHT, *Numerical optimization*, Springer, New York, 1999.
 - [41] A. POLYAKOV, *Nonlinear feedback design for fixed-time stabilization of linear control systems*, IEEE Transactions on Automatic Control, 57 (2012), pp. 2106–2110.
 - [42] M. RAISSI, P. PERDIKARIS, AND G. E. KARNIADAKIS, *Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations*, Journal of Computational physics, 378 (2019), pp. 686–707.
 - [43] M. F. REIS, A. P. AGUIAR, AND P. TABUADA, *Control barrier function-based quadratic programs introduce undesirable asymptotically stable equilibria*, IEEE Control Systems Letters, 5 (2020), pp. 731–736.
 - [44] W. REN, R. M. JUNGERS, AND D. V. DIMAROGONAS, *Razumikhin and krasovskii approaches for safe stabilization*, Automatica, 146 (2022), p. 110563.
 - [45] M. Z. ROMDLONY AND B. JAYAWARDHANA, *Stabilization with guaranteed safety using control lyapunov-barrier function*, Automatica, 66 (2016), pp. 39–47.
 - [46] J. D. SÁNCHEZ-TORRES, D. GÓMEZ-GUTIÉRREZ, E. LÓPEZ, AND A. G. LOUKIANOV, *A class of predefined-time stable dynamical systems*, IMA Journal of Mathematical Control and Information, 35 (2018), pp. i1–i29.
 - [47] E. D. SONTAG, *A ‘universal’ construction of artstein’s theorem on nonlinear stabilization*, Systems & control letters, 13 (1989), pp. 117–123.
 - [48] K. G. VAMVOUDAKIS AND N.-M. T. KOKOLAKIS, *Synchronous Reinforcement Learning-Based Control for Cognitive Autonomy*, vol. 8, Now Publishers, Boston - Delft, 2020.
 - [49] A. J. VAN DER SCHAFT, *Nonlinear State Space H_∞ Control Theory*, Springer, Berlin, 1993.
 - [50] P. WIELAND AND F. ALLGÖWER, *Constructive safety using control barrier functions*, IFAC Proceedings Volumes, 40 (2007), pp. 462–467.
 - [51] W. XIAO, C. G. CASSANDRAS, AND C. BELTA, *Safe Autonomy with Control Barrier Functions: Theory and Applications*, Springer Nature, Cham, 2023.
 - [52] X. XU, P. TABUADA, J. W. GRIZZLE, AND A. D. AMES, *Robustness of control barrier functions for safety critical control*, IFAC-PapersOnLine, 48 (2015), pp. 54–61.
 - [53] C. ZHU, R. H. BYRD, P. LU, AND J. NOCEDAL, *Algorithm 778: L-bfgs-b: Fortran subroutines for large-scale bound-constrained optimization*, ACM Transactions on mathematical software (TOMS), 23 (1997), pp. 550–560.