

Unified Multi-Rate Control: from Low-Level Actuation to High-Level Planning

Ugo Rosolia, *Member, IEEE*, Andrew Singletary, *Member, IEEE*, and Aaron D. Ames, *Fellow, IEEE*

Abstract—In this paper we present a hierarchical multi-rate control architecture for nonlinear autonomous systems operating in partially observable environments. Control objectives are expressed using syntactically co-safe Linear Temporal Logic (LTL) specifications and the nonlinear system is subject to state and input constraints. At the highest level of abstraction, we model the system-environment interaction using a discrete Mixed Observable Markov Decision Process (MOMDP), where the environment states are partially observed. The high-level control policy is used to update the constraint sets and cost function of a Model Predictive Controller (MPC) which plans a reference trajectory. Afterwards, the MPC planned trajectory is fed to a low-level high-frequency tracking controller, which leverages Control Barrier Functions (CBFs) to guarantee bounded tracking errors. Our strategy is based on model abstractions of increasing complexity and layers running at different frequencies. We show that the proposed hierarchical multi-rate control architecture maximizes the probability of satisfying the high-level specifications while guaranteeing state and input constraint satisfaction. Finally, we tested the proposed strategy in simulations and experiments on examples inspired by the Mars exploration mission, where only partial environment observations are available.

Index Terms—partially observable, noisy observations, predictive control, control barrier function, multi-rate control, hierarchical control.

I. INTRODUCTION

Control design for complex cyber-physical systems, which are described by continuous and discrete variables, is usually divided into different layers [1]–[11]. Each layer is designed using model of increasing accuracy and complexity, which allow the controller to take high-level decisions—e.g., perform an overtaking maneuver—and to compute low-level commands—e.g., the input current to a motor. High-level decisions and low-level control actions are computed at different frequencies and the interaction between layers should be taken into account to guarantee safety of the closed-loop system [2].

In this work, we present a multi-rate hierarchical control scheme for nonlinear systems operating in partially observable environments. Our architecture, which is composed by three layers running at different frequencies, guarantees constraint satisfaction and maximization of the closed-loop probability of

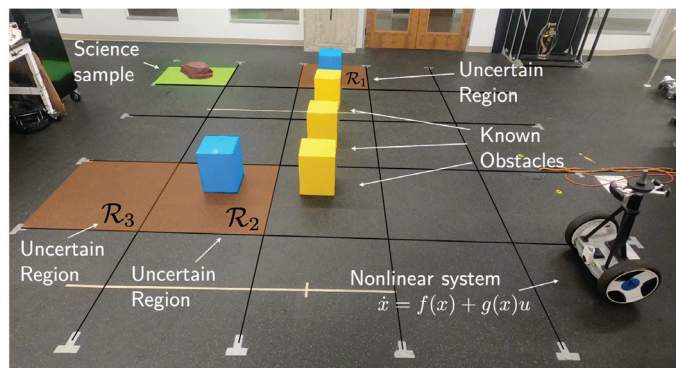


Fig. 1. This figure shows an environment composed of 25 cells, 3 obstacles (yellow and blue boxes) and 3 uncertain regions (light brown). In this example the goal of the controller is to explore the state space in order to find a science sample.

satisfying the high-level specifications. At the lowest level, we leverage the continuous time nonlinear system model to guarantee a bounded tracking error. The mid-level planning layer computes a reference trajectory using a simplified prediction model and the low-level tracking error bounds. Finally, at the highest level of abstraction we model the system-environment interaction using Mixed Observable Markov Decision Processes (MOMDPs), which allows us to account for partial environment observations.

A. Related Work

Control policies for high-level decision making are usually synthesized using discrete model abstractions. The high-level control objectives are often expressed by Linear Temporal Logic (LTL) formulas [12], as they are a formalism to express high-level system behaviors using logical and temporal operators [12]. Motion planning with LTL and syntactically co-safe LTL (scLTL) specifications has been widely studied in literature [2]–[4], [6], [7], [13]–[22]. For deterministic systems with finite-state spaces several approaches and toolboxes are available for synthesis [2]–[4], [13]–[15]. When the system-environment interaction are uncertain, the high-level abstractions are described by discrete Markov Decision Processes (MDPs) and the high-level decision making problem can be solved exactly using dynamic programming, policy iteration, and linear programming strategies [23]. On the other hand, when the system dynamics are uncertain and only partial observations are available, the system-environment interaction

U. Rosolia, A. Singletary and A. D. Ames are with the AMBER lab at the California Institute of Technology, Pasadena, CA, USA, e-mail: {urosolia, asinglet, ames}@caltech.edu. The authors would like to acknowledge the support by the National Science Foundation award #1932091.

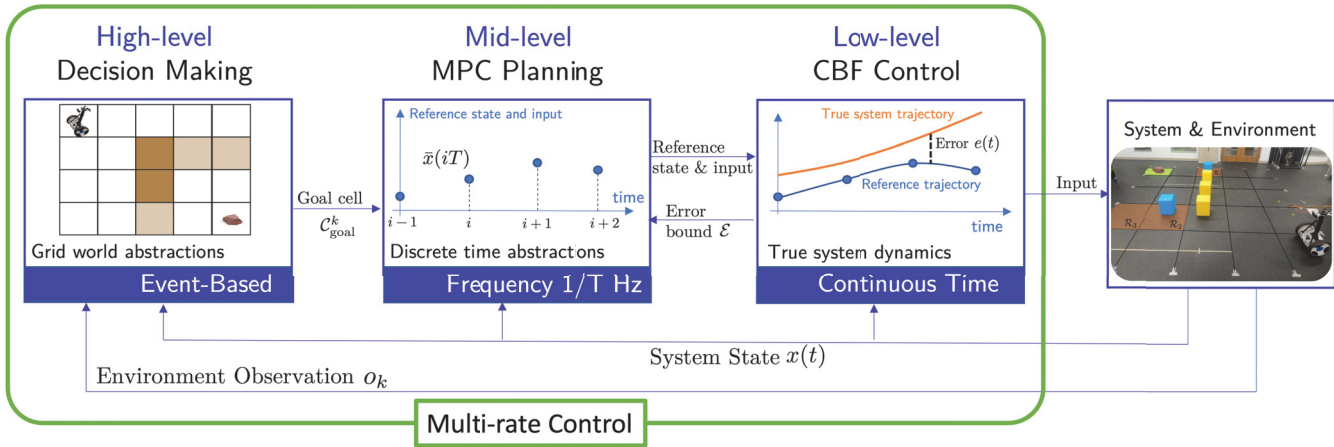


Fig. 2. Multi-rate control architecture. The high-level decision maker leverages the system's state $x(t)$ and partial environment observations o_k to compute a goal cell, the constraint set and the goal positions, which are fed to the mid-level MPC planner. The planner computes a reference trajectory given the tracking error bounds $\mathcal{E}$ from the low-level tracking controller. Finally at the lowest level, the control action is computed summing up the mid-level input $u_m(t)$ and the low-level input $u_l(t)$.

can be modeled using discrete Partially Observable Markov Decision Processes (POMDPs). Computing a control policy in POMDPs settings is NP-hard [24], but approximate solutions can be computed using finite state controllers [25] and performing point-based approximations [26].

Given a high-level decision, reachability-based techniques [2], [3] or simulation-based abstractions [4], [5] may be used to compute a goal set for the continuous time system, e.g., a subset of a lane where we would like to drive the vehicle when performing an overtaking maneuver. Therefore, the input to the system's actuators is computed solving mid-level planning and low-level control problems, which have been studied extensively in literature [10], [27]–[38]. The planning problem is usually defined for a simplified model and the resulting reference trajectory is then tracked using low-level controllers, which leverage the nonlinear system dynamics. Tracking controllers may be synthesized using Hamilton-Jacobi (HJ) reachability analysis [10] or sum-of-squares programming [31], [32]. Another strategy to solve mid-level planning and low-level control problems is to use nonlinear tube MPC [33]–[36], [39], where the difference between the planned trajectory and the actual one is over approximated using Lyapunov based analysis or Lipschitz properties of the nonlinear dynamics. When the planned trajectory is computed without taking into account tracking errors, safety can be guaranteed using filters which, given a desired mid-level command, compute a safe control action using CBFs [27]–[29], feasibility of an MPC problem [30], or reachability analysis [9]. A different strategy that can be used to bridge the gap between high-level decision making and low-level control is to leverage CBFs [37], [38]. However, these strategies compute control actions without forecasting the evolution of the system's trajectory and they may result in sub-optimal behaviors.

As discussed next, our approach leverages both MPC and CBF policies to compute low-level commands given high-level decisions. The forecast from the MPC planning layers is used to compute a feedforward term that allows us to mitigate the

myopic nature of CFBs, which are used to guarantee safety at the continuous time layer [28].

B. Contribution

Our contribution is threefold. First, we introduce a mid-level planner that leverages two MPC problems with time-varying constraint sets and cost functions. These time-varying components are given by the high-level decision maker and they can jeopardize the feasibility of standard MPC schemes, which are designed to steer the system to a time-invariant goal state. For instance, the safety guarantees from [40]–[41] are lost when the goal state and constraints are updated online during the execution of the control task. To overcome limitations of standard time-invariant approaches, we propose a contingency scheme where at each time step we solve at most two MPC problems. This strategy guarantees feasibility of the planner with time-varying components. In particular, in Algorithm 2 we introduce a contingency MPC problem that is defined by updating the time-varying components as a function of the latest planned optimal trajectory.

Second, we show how to integrate a CLF-CBF tracking controller with a mid-level planner. We present a constraint tightening approach that accounts for the low-level tracking error and we demonstrate that the resulting multi-rate control architecture guarantee safety, when a local reachability assumption on the system dynamics is satisfied. Such reachability assumption, which is tailored to navigation problems, together with the proposed contingency scheme allows us to avoid the construction of finite state abstractions defined over the entire state space. Furthermore compared to the constraint tightening from our previous work [40], the proposed constraint tightening builds upon ideas from the fixed-tube robust MPC strategy [41], where the initial state of the planned trajectory is an optimization variable. For this reason, the proposed constraint tightening does not require the online computation of robust reachable sets to formulate the MPC problem and therefore it is computationally more efficient than the formulation from [40].

Third, we show how to model the system-environment interaction using Mixed Observable Markov Decision Processes (MOMDPs), where the system's state is fully observable and the environment's state is partially observable. We build upon [42] that introduced the synthesis process for systems with discrete state and action spaces, and we focus on the multi-layer hierarchical control design for systems with continuous states and actions. In particular, we show how to leverage high-level decisions from the MOMDP to construct the MPC time-varying components, and we demonstrate that the proposed multi-layer hierarchical control strategy maximizes the probability of satisfying the high-level specifications. Finally, we test our strategy on navigation tasks as the one shown in Figure 1, where a Segway like-robot has to find science samples while navigating a partially observable environment.

This paper is organized as follows. The background material is discussed in Section II. Section III describes the problem under study. First, we introduce the system and environment models, and afterwards the control design objectives. The hierarchical architecture is introduced in Section IV, where we present the high-level decision maker, the mid-level MPC planer, and the low-level CBF controller. The closed-loop properties are discussed in Section V. Finally, we illustrate the effectiveness of the proposed strategy with high-fidelity simulations and hardware experiments.

II. PRELIMINARIES

Notation: The Minkowski sum of two sets $\mathcal{X} \subset \mathbb{R}^{n_x}$ and $\mathcal{Y} \subset \mathbb{R}^{n_x}$ is denoted as $\mathcal{X} \oplus \mathcal{Y}$, and the Pontryagin difference as $\mathcal{X} \ominus \mathcal{Y}$. $\mathcal{K}^e$ is the set of extended class- $\mathcal{K}^e$ functions β which are strictly increasing and $\beta(0) = 0$. For a set $\mathcal{A} \subset \mathbb{R}^{n_x}$ and a vector $x \in \mathbb{R}^{n_x}$, we denote the projection

$$\text{Proj}(x, \mathcal{A}) = \underset{d \in \mathcal{A}}{\text{argmin}} \|x - d\|_2,$$

and the cardinality of the set $\mathcal{A}$ as $|\mathcal{A}|$. We define $\mathbb{Z}_{0+} = \{0, 1, 2, \dots\}$ and $\mathbb{R}_{0+} = \{x \in \mathbb{R}^{n_x} | x \geq 0\}$ which denote the set of positive integers and real numbers, respectively. Finally, given $t \in \mathbb{R}_{0+}$ and $T \in \mathbb{Z}_{0+}$ we define $\lfloor t/T \rfloor = \text{floor}(t/T)$.

Specifications: High-level objectives are expressed using syntactically co-safe Linear Temporal Logic (scLTL) specifications. For a set of atomic proposition $\mathcal{AP}$, an scLTL specification is defined as follows:

$$\psi := p \mid \neg p \mid \psi_1 \wedge \psi_2 \mid \psi_1 \vee \psi_2 \mid \psi_1 U \psi_2 \mid \bigcirc \psi,$$

where the atomic proposition $p \in \mathcal{AP}$ and ψ, ψ_1, ψ_2 are scLTL formulas, which can be defined using the logic operators negation ($\neg$), conjunction ($\wedge$) and disjunction ($\vee$). Furthermore, scLTL formulas can be specified using the temporal operators until (U) and next ($\bigcirc$). Each atomic proposition p is associated with a subset of the high-level state space $\mathcal{P}$ and a high-level state ω_k satisfies the proposition p if $\omega_k \in \mathcal{P}$. Finally, satisfaction of a specification ψ for the trajectory $\omega_k = [\omega_k, \omega_{k+1}, \dots]$, denoted by

$$\omega_k \models \psi, \quad (1)$$

is recursively defined as follows: *i*) $\omega_k \models p \iff \omega_k \in \mathcal{P}$, *ii*) $\omega_k \models \psi_1 \wedge \psi_2 \iff (\omega_k \models \psi_1) \wedge (\omega_k \models \psi_2)$, *iii*)

$\omega_k \models \psi_1 \vee \psi_2 \iff (\omega_k \models \psi_1) \vee (\omega_k \models \psi_2)$, *iv*) $\omega_k \models \psi_1 U \psi_2 \iff \omega_l \models \psi_2$ and $\omega_j \models \psi_1, \forall j \in \{k, \dots, l-1\}$, *v*) $\omega_k \models \bigcirc \psi \iff \omega_{k+1} \models \psi$. Please refer to [43, Chapter 3] for further details.

III. PROBLEM FORMULATION

This section describes the problem formulation. First, we introduce the continuous system dynamics. Afterwards, we present the discrete environment model. Finally, we describe the synthesis goals and we summarize the overall control architecture from Figure 2.

System Model: As discussed in the introduction, our goal is to design a controller for nonlinear dynamical systems. In particular, we consider nonlinear control affine systems of the following form:

$$\dot{x}(t) = f(x(t)) + g(x(t))u(t), \quad (2)$$

where f and g are Lipschitz continuous, the input $u(t) \in \mathbb{R}^{n_u}$ and the state $x(t) = [p^\top(t), q^\top(t)]^\top \in \mathbb{R}^{n_x}$ for the position vector $p(t) \in \mathbb{R}^{n_p}$ and the vector $q(t) \in \mathbb{R}^{n_q}$ collecting the remaining states. Furthermore, the above system is subject to the following state and input constraints:

$$u(t) \in \mathcal{U}, p(t_i) \in \mathcal{X}_p \text{ and } q(t_i) \in \mathcal{X}_q, \quad (3)$$

for all $t \in \mathbb{R}_{0+}$ and $t_i = iT$ for all $i \in \mathbb{Z}_{0+}$. The time constant T is specified by the user and, as it will be clear later on, it defines the frequency at which the controller updates the planned trajectory. In the above equation (3), $\mathcal{X}_p$ represents the free space and $\mathcal{X}_q$ is a user-defined constraint set.

Remark 1. We consider state constraints which are enforced pointwise in time to streamline the presentation. The proposed control strategy can be extended to account for constraints which must hold for all time $t \in \mathbb{R}_{0+}$. In this case, it is required to modify the low-level controller as discussed in [40].

Environment Model: We consider nonlinear dynamical systems operating in partially observable environments, which are partitioned into $\mathcal{C}_1, \dots, \mathcal{C}_c$ cells as in the example from Figure 1. We assume that the state of the system is perfectly observable, but we are given only partial observations about the environment state. Thus, at the highest level of abstraction, we model the interaction between the nonlinear system (2) and the environment using a Mixed Observable Markov Decision Process (MOMDP). A MOMDP provides a sequential decision-making formalism for high-level planning under mixed full and partial observations [44] and it is defined as a tuple $(\mathcal{S}, \mathcal{Z}, \mathcal{A}, \mathcal{O}, T_s, T_z, O)$, where

- $\mathcal{S} = \{1, \dots, |\mathcal{S}|\}$ is a set of fully observable states;
- $\mathcal{Z} = \{1, \dots, |\mathcal{Z}|\}$ is a set of partially observable states;
- $\mathcal{A} = \{1, \dots, |\mathcal{A}|\}$ is a set of actions;
- $\mathcal{O} = \{1, \dots, |\mathcal{O}|\}$ is the set of observations for the partially observable state $z \in \mathcal{Z}$;
- The indicator function¹ $T_s : \mathcal{S} \times \mathcal{Z} \times \mathcal{A} \times \mathcal{S} \rightarrow \{0, 1\}$ equals one if the system will transition to a state s' given

¹We introduced the indicator function as it will be used later on to compute the belief vector update.

the action a and current state (s, z) , i.e.,

$$T_s(s, z, a, s') = \begin{cases} 1 & \text{If } s' = f_s(s, z, a) \\ 0 & \text{Else} \end{cases},$$

where the high-level update function $f_s : \mathcal{S} \times \mathcal{Z} \times \mathcal{A} \rightarrow \mathcal{S}$.

- The function $T_z : \mathcal{S} \times \mathcal{Z} \times \mathcal{A} \times \mathcal{S} \times \mathcal{Z} \rightarrow [0, 1]$ describes the probability of transitioning to a state z' given the action a , the successor observable state s' , and the system's current state (s, z) , i.e.,

$$T_z(s, z, a, s', z') := P(z_{k+1} = z' | s_k = s, z_k = z, a_k = a, s_{k+1} = s');$$

- The function $O : \mathcal{S} \times \mathcal{Z} \times \mathcal{A} \times \mathcal{O} \rightarrow [0, 1]$ describes the probability of observing the measurement $o \in \mathcal{O}$, given the current state of the system $(s', z') \in \mathcal{S} \times \mathcal{Z}$ and the action a applied at the previous time step, i.e.,

$$O(s', z', a, o) := P(o_k = o | s_k = s', z_k = z', a_{k-1} = a);$$

MOMDPs were introduced in [44] to model systems where a subspace of the state space is perfectly observable². In this work, the high-level observable state s represents the location of the system, i.e., the grid cell containing the position vector $p(t)$ which is part of state $x(t) = [p^\top(t), q^\top(t)]^\top$ of the nonlinear system (2). On the other hand, the definition of the partially observable state z depends on the application, and it describes how the environment may affect the evolution of the system. For example, it may be used to model external events (e.g., rain, wind, etc) that would affect the traversability of specific regions of the state space. The evolution of the environment state z may be stochastic and, most importantly, it is not perfectly observable. Thus, the controller has to make decisions based on the belief about the environment. For example, when the objective is to reach a goal location before a deadline and only partial knowledge about the traversability of the terrain is given, the controller should follow a path that maximizes the probability of reaching the goal in time, given our belief about the environment.

More formally, control actions are computed based on the environment belief vector $b_k \in \mathcal{B} = \{b \in \mathbb{R}^{|\mathcal{Z}|} : \sum_{z=1}^{|\mathcal{Z}|} b(z) = 1\}$ representing the posterior probability that the partially observable state environment z_k equals $z \in \mathcal{Z}$, i.e., $b_k = [b_k^{(1)}, \dots, b_k^{(|\mathcal{Z}|)}]$ with

$$b_k^{(z)} := \mathbb{P}(z_k = z | o_k, s_k, a_{k-1}), \quad \forall z \in \{1, \dots, |\mathcal{Z}|\}.$$

In the above definition, at time k the observation vector $o_k = [o_0, \dots, o_k]$, the observable state vector $s_k = [s_0, \dots, s_k]$, and the action vector $a_{k-1} = [a_0, \dots, a_{k-1}]$. Notice that the evolution of the environment belief vector b_k is stochastic as it is a function of the noisy observation vector $o_k = [o_0, \dots, o_k]$. Therefore, the planned path that maximizes the probability of completing the task should be computed online at after collecting the observation o_k about the environment's state and updating the belief vector b_k .

²We introduced a special case of the MOMDP from [44] where the transition function of the observable state is deterministic given both the current observable and unobservable states.

Synthesis Objectives: Given the system's state $x(t) \in \mathbb{R}^{n_x}$ and k observations $o_{k-1} = [o_0, \dots, o_{k-1}] \in \mathcal{O}^k$ about the environment, our goal is to design a control policy

$$\pi : \mathbb{R}^{n_x} \times \mathcal{O}^k \rightarrow \mathcal{U}, \quad (4)$$

which maps the state $x(t)$ and the observation vector o_{k-1} to the continuous control action $u \in \mathcal{U}$. Furthermore, the control policy (4) should guarantee that state and input constraints (3) are satisfied and that the probability of satisfying the specification (1) is maximized. Notice that standard control strategies for nonlinear systems can be used to guarantee constraint satisfaction [27]–[29], [33]–[36], [39]. Furthermore, standard decision making methodologies for Partially Observable Markov Decision Processes (POMDPs) can be used to synthesize a control policy which maximizes the probability of satisfying the specification [6], [7], [17]–[20]. In this paper, we bridge the gap between the two communities and we propose a hierarchical control scheme for nonlinear systems operating in partially observable environments, which guarantees that state and input constraints are satisfied and that the probability of satisfying the specifications is maximized.

Navigation Example: Figure 1 shows our motivating example, where a Segway has to reach a goal cell while avoiding known obstacles and exploring uncertain regions, which may be traversable with some probability. The Segway dynamics are nonlinear and the system is open-loop unstable, for this reason it is required a low-level high frequency controller that stabilizes the system during operations. On the other hand, at the highest level of abstraction we model the system using the discrete state $s_k \in \mathcal{S}$, which denotes the grid cell containing the nonlinear system (2), and the environment state $z_k \in \mathcal{Z}$ representing the traversability of the uncertain regions $\mathcal{R}_1, \mathcal{R}_2$ and $\mathcal{R}_3$. In this navigation example, where the traversability of n_o regions is unknown, we define the vector z_k as follows:

$$z_k = [z_k^{(1)}, \dots, z_k^{(n_o)}] \in \mathcal{Z},$$

where each entry $z_k^{(i)}$ equals one if the i -th region is traversable and zero otherwise. For instance in the settings from Figure 1, the environment's state $z_k = [0, 0, 1]$ as regions $\mathcal{R}_1$ and $\mathcal{R}_2$ are not traversable and region $\mathcal{R}_3$ is traversable.

Strategy Overview: We summarize the proposed multi-rate control architecture depicted in Figure 2. The key idea is to divide the controller into three layers and compute the control action $u(t)$ as the summation of a high-frequency component $u_l(t)$ and a low-frequency component $u_m(t)$, i.e.,

$$u(t) = u_l(t) + u_m(t).$$

At the lowest level, the control action u_l is updated continuously (at high frequency) and it is computed using Control Barrier Functions (CFBs), which leverage the full-nonlinear model (2) to track a *reference trajectory* $\bar{x}(t)$. The middle layer updates at a constant frequency the reference trajectory $\bar{x}(t)$ and reference input u_m , which is computed using a Model Predictive Controller (MPC). This reference trajectory steers the system from the current state $x(t)$ to a *goal cell* $\mathcal{C}_{\text{goal}}^k$. Finally, the high-level planner computes the goal cell $\mathcal{C}_{\text{goal}}^k$ based on partial observations o_k about the environment.

IV. UNIFIED MULTI-RATE ARCHITECTURE

In this section, we describe the multi-rate control architecture. First, we design a low-level CLF-CBF controller, which tracks a reference state-input trajectory and guarantees bounded tracking errors. Afterwards, we show how to update the state-input reference trajectory leveraging an MPC, which is designed using a goal state computed from a discrete high-level decision maker. Finally, we introduce the hierarchical multi-rate architecture, which guarantees that the synthesis objectives from Section III are satisfied.

A. Low-Level Control

We leverage CBFs and CLFs to design a low-level tracking controller for the nonlinear system (2). CBFs guarantee safety for nonlinear system [28], but they are suboptimal as the control action is computed without forecasting the system's trajectory. For this reason, we use CBFs to enforce safety around a reference state-input trajectory that is computed at low frequency by the mid-level planner, as shown in Figure 2.

Error Model: At the lowest layer, the goal of the controller is to track a reference trajectory $\bar{x}(t)$. We assume that the reference trajectory is given by the following Linear Time-Varying (LTV) model:

$$\Sigma_{\bar{x}} : \begin{cases} \dot{\bar{x}}(t) = A_{[t/T]} \bar{x}(t) + B_{[t/T]} u_m(t), & t \in \mathcal{T} \\ \bar{x}^+(t) = \Delta_{\bar{x}}(x(t)), & t \in \mathcal{T}^c \end{cases}, \quad (5)$$

where $\mathcal{T}^c = \cup_{j=0}^{\infty} \{jT\}$, $\mathcal{T} = \cup_{j=0}^{\infty} (jT, (j+1)T)$ and the time T from (3) is specified by the user. Furthermore, we denote $\bar{x}^-(t) = \lim_{\tau \nearrow t} \bar{x}(\tau)$ and $\bar{x}^+(t) = \lim_{\tau \searrow t} \bar{x}(\tau)$ as the right and left limits of the reference trajectory $\bar{x}(t) \in \mathbb{R}^{n_x}$, which is assumed right continuous. In the above system, the reference input $u_m(t) \in \mathbb{R}^{n_u}$ and the reset map $\Delta_{\bar{x}} : \mathbb{R}^{n_x} \rightarrow \mathbb{R}^{n_x}$ maps the current state of the system $x(t)$ to the state $\bar{x}(t)$ of the reference trajectory. Both the reference input and the reset map are given by the middle layer as we will discuss in Section IV-B. Finally, the time-varying matrices $(A_{[t/T]}, B_{[t/T]})$ are known and, in practice, may be computed linearizing the system dynamics (2), as discussed in the result section.

Given the nonlinear system (2) and the LTV model (5), we define the error state $e(t) = x(t) - \bar{x}(t)$ and the associated error dynamics:

$$\Sigma_e : \begin{cases} \dot{e}(t) = f_e(x(t), \bar{x}(t), u_l(t) + u_m(t), t), & t \in \mathcal{T} \\ e^+(t) = x^+(t) - \bar{x}^+(t), & t \in \mathcal{T}^c \end{cases} \quad (6)$$

where the time-varying error dynamics are:

$$\begin{aligned} f_e(x, \bar{x}, u_l + u_m, t) \\ = f(x) + g(x)(u_l + u_m) - (A_{[t/T]} \bar{x} + B_{[t/T]} u_m). \end{aligned}$$

In the above definition, we dropped the dependence on time for states and inputs to simplify the notation. Furthermore, we introduce the low-level input constraint set $\mathcal{U}_l \subset \mathcal{U}$ and the mid-level input constraint set $\mathcal{U}_m \subset \mathcal{U}$ which partition the input space, i.e.,

$$\mathcal{U}_l \oplus \mathcal{U}_m = \mathcal{U}.$$

Next, we design a low-level controller which guarantees that the reference trajectory $\bar{x}(t)$ from the LTV model (5) is tracked within some error bounds.

Control Barrier and Lyapunov Functions: We show how to design a tracking controller using CBFs and CLFs [28]. First, we define the candidate Lyapunov function

$$V(e) = \|e\|_Q, \quad (7)$$

where $\|e\|_Q = e^\top Q e$. Furthermore, we introduce the following safe set for the error dynamics (6):

$$\mathcal{E} = \{e \in \mathbb{R}^{n_x} : h_e(e) \geq 0\} \subset \mathbb{R}^{n_x}. \quad (8)$$

The above function h_e is defined by the user and it depends on the application as discussed in the result section.

Finally, the CBF associated with the safe set (8), and the CLF from (7) are used to define the following CLF-CBF Quadratic Program (QP):

$$\begin{aligned} v_l^*(t) = \operatorname{argmin}_{v_l \in \mathcal{U}_l, \gamma} & \|v_l\|_2 + c_1 \gamma^2 \\ \text{subject to} & \frac{\partial V(e)}{\partial e} f_e(x, \bar{x}, v_l + u_m) \leq -c_2 V(e) + \gamma \\ & \frac{\partial h_e(e)}{\partial e} f_e(x, \bar{x}, v_l + u_m) \geq -\alpha(h_e(e)), \end{aligned} \quad (9)$$

where we dropped the time dependence to simplify the notation, and $v_l \in \mathcal{U}_l$ is the low-level control action. In the above QP, the parameters $c_1 \in \mathbb{R}_{0+}$, $c_2 \in \mathbb{R}_{0+}$, and $\alpha \in \mathcal{K}^e$. Given the optimal control action $v_l^*(t)$ from the QP (9), the low-level control policy is defined as follows:

$$u_l(t) = \pi_l(x(t), \bar{x}(t), u_m(t)) = v_l^*(t). \quad (10)$$

Assumption 1. The CLF-CBF QP (9) is feasible for all $e = x - \bar{x} \in \mathcal{E}$ and for all $u_m \in \mathcal{U}_m$.

Remark 2. We underline that Assumption 1 is satisfied for some $\alpha \in \mathcal{K}^e$ when the set $\mathcal{E}$ is a robust control invariant for system (6) with $u_m(t) \in \mathcal{U}_m$ and mild assumptions on the Lie derivative of (6) hold (see [28] for further details). The set $\mathcal{E}$ may be hard to compute and standard techniques are based on HJB reachability analysis [10], SOS programming [32], Lyapunov-based methods [33], and Lipschitz properties of the system dynamics [36], [45].

The low-level control policy (10) guarantees that the difference between the evolution of the nonlinear system (2) and the LTV model (5) is bounded. Indeed, when Assumption 1 is satisfied, the CLF-CBF QP (9) guarantees invariance of the safe set (8) for all $t \in (iT, (i+1)T)$ and $i \in \mathbb{Z}_{0+}$, as discussed in Section V. Next, we show how to design a mid-level planner which leverages the safe set $\mathcal{E}$ from (8).

B. Mid-Level Planning

In this section we describe the mid-level planning strategy. At this level of abstraction, we assume that we are given a goal grid cell where we would like to steer the system. Afterwards, we compute a reference state-input trajectory using an MPC, which leverages a simplified model and the tracking error bounds from the previous section.

Grid Model: Given the state $x(t) = [p^\top(t), q^\top(t)]^\top$, we define the current grid cell $\mathcal{C}_{\text{curr}}^k$, which contains the nonlinear system (2) for time $t \in [t^k, t^{k+1})$, i.e.,

$$p(t) \in \mathcal{C}_{\text{curr}}^k \subset \mathcal{X}_p, \forall t \in [t^k, t^{k+1}). \quad (11)$$

Similarly, we define the goal cell $\mathcal{C}_{\text{goal}}^k$, which represents the region where we want to steer the system for time $t \in [t^k, t^{k+1})$. Finally, we introduce the goal equilibrium sets $\mathcal{X}_{\text{curr}}^k$ and $\mathcal{X}_{\text{goal}}^k$, which collect the unforced equilibrium states that are contained into $\mathcal{C}_{\text{curr}}^k$ and $\mathcal{C}_{\text{goal}}^k$, i.e., for $i \in \{\text{curr}, \text{goal}\}$

$$\mathcal{X}_i^k = \{x = [p, q] \in \mathbb{R}^{n_x} | p \in \mathcal{C}_i^k, \dot{x} = f(x) = 0\} \subset \mathbb{R}^{n_x}. \quad (12)$$

Throughout this section, we assume that t^k , $\mathcal{X}_{\text{goal}}^k$, $\mathcal{C}_{\text{curr}}^k$, and $\mathcal{C}_{\text{goal}}^k$ are given by the high-level planner and we synthesize a controller to drive the system from the current cell $\mathcal{C}_{\text{curr}}^k$ to the goal cell $\mathcal{C}_{\text{goal}}^k$.

Model Predictive Control: We design an MPC to compute the mid-level input $u_m(t)$ that defines the evolution of the reference trajectory (5) and to define the reset map $\Delta_{\bar{x}}$ for the LTV model (5). The MPC problem is solved at $1/T$ Hz and therefore the reference mid-level control input is piecewise constant, i.e.,

$$\dot{u}_m(t) = 0 \quad \forall t \in \mathcal{T} = \bigcup_{k=0}^{\infty} (kT, (k+1)T).$$

Next, we introduce the following discrete time linear model:

$$\bar{x}^d((i+1)T) = \bar{A}_i \bar{x}^d(iT) + \bar{B}_i u(iT), \quad (13)$$

where for all $i \in \mathbb{Z}_{0+}$

$$\bar{A}_i = e^{A_{[iT/T, T]} T} \text{ and } \bar{B}_i = \int_0^T e^{A_{[iT/T, T]} (T-\eta)} B_{[iT/T, T]} d\eta,$$

for the matrices $A_{[iT/T, T]}$ and $B_{[iT/T, T]}$ defined in (5). Now notice that, as the mid-level input u_m is piecewise constant, if at time $t_i = iT$ the state of the nominal model (5) $\bar{x}(iT) = \bar{x}^d(iT)$, then at time $t_{i+1} = (i+1)T$ we have that

$$\bar{x}^-((i+1)T) = \bar{x}^d((i+1)T). \quad (14)$$

Given the discrete time model (13), at time $t_i = iT \in \mathcal{T}^c$ we solve the following finite time optimal control problem:

$$J(x(iT), N) =$$

$$\min_{v_t, x_{t|i}^d} \|x_{t|i}^d - x(iT)\|_{Q_e} + \sum_{t=i}^{i+N-1} \ell(x_{t|i}^d, v_{t|i}) + \|p_{i+N|i}^d - p_{\text{goal}}^k\|_{Q_f}$$

$$\begin{aligned} \text{subject to } & x_{t+1|i}^d = \bar{A}_t x_{t|i}^d + \bar{B}_t v_{t|i}^d \\ & x_{t|i}^d = \begin{bmatrix} p_{t+1|i}^d \\ q_{t+1|i}^d \end{bmatrix} \in \mathcal{X}_{p,q}^k \ominus \mathcal{E}, \quad v_{t|i}^d \in \mathcal{U}_m \\ & x_{i|i}^d - x(iT) \in \mathcal{E} \\ & x_{i+N|i}^d \in \mathcal{X}_{\text{goal}}^k \ominus \mathcal{E}_p, \forall t = \{i, \dots, i+N-1\} \end{aligned} \quad (15)$$

where $\mathcal{E}$ is defined in (8), $\|p\|_Q = p^\top Q p$,

$$\mathcal{X}_{p,q}^k = \left\{ x = \begin{bmatrix} p \\ q \end{bmatrix} \in \mathbb{R}^{n_x} | p \in \mathcal{C}_{\text{curr}}^k \cup \mathcal{C}_{\text{goal}}^k \text{ and } q \in \mathcal{X}_q \right\} \quad (16)$$

and

$$\mathcal{E}_p = \left\{ e = \begin{bmatrix} e_p \\ 0 \end{bmatrix} \in \mathbb{R}^{n_x} | \exists e_q \in \mathbb{R}^{n_p} \text{ and } \begin{bmatrix} e_p \\ e_q \end{bmatrix} \in \mathcal{E} \right\}. \quad (17)$$

Notice that the MPC problem (15) is designed based on the time-varying components $\mathcal{X}_{\text{goal}}^k$, $\mathcal{C}_{\text{curr}}^k$, $\mathcal{C}_{\text{goal}}^k$, p_{goal}^k which are given by the high-level decision maker, as shown in Figure 2. Problem (15) computes a sequence of open-loop actions $v_t^d = [v_{t|t}^d, \dots, v_{t+N|t}^d]$ and an initial condition $x_{i|i}^d$ such that the predicted trajectory steers the system to the terminal set $\mathcal{X}_{\text{goal}}^k$, while minimizing the cost and satisfying state and input constraints. Let $v_{t|t}^{d,*} = [v_{t|t}^{d,*}, \dots, v_{t+N|t}^{d,*}]$ be the optimal solution and $[x_{t|t}^{d,*}, \dots, x_{t+N|t}^{d,*}]$ the associated optimal trajectory, then the mid-level policy is

$$\Pi_m : \begin{cases} u_m(t) = \pi_m(x(t), N) = v_{t|t}^{d,*} & t \in \mathcal{T}^c \\ \dot{u}_m(t) = 0 & t \in \mathcal{T} \end{cases}. \quad (18)$$

Finally, we define the reset map from the LTV model (5) as:

$$\Delta_{\bar{x}}(x(t)) = x_{t|t}^{d,*}. \quad (19)$$

Assumption 2. Consider the equilibrium set $\mathcal{X}_{\text{curr}}^k$ defined in equation (12). For all states $x(t) \in \mathcal{X}_{\text{curr}}^k \oplus \mathcal{E}$ problem (15) is feasible with horizon N .

The above assumption is satisfied when any equilibrium state $\bar{x} \in \mathcal{X}_{\text{curr}}^k$ of the discrete time system (13) can be steered to the goal equilibrium set $\mathcal{X}_{\text{goal}}^k$ in at most N time steps. More formally, Assumption 2 holds when, for the discrete time system (13), $\mathcal{X}_{\text{goal}}^k$ is N -step backward reachable from the set $\mathcal{X}_{\text{curr}}^k$.

In Section V, we will show that when the nonlinear system (2) and the LTV system (5) are in closed-loop with the low-level policy (10) and the mid-level policy (18), then state and input constraints (3) are satisfied for system (2). Furthermore, the nonlinear system (2) is steered from the current cell $\mathcal{C}_{\text{curr}}^k$ to the goal cell $\mathcal{C}_{\text{goal}}^k$ in finite time.

Remark 3. We highlight that also RRT-based methods can be combined with CLF-CBF to design a multi-rate control architecture. In particular, it would be possible to leverage sampling-based methods, such as [46]–[49], to repeatedly solve online problem (15). Notice that it is important to consider the constraint tightening from problem (15) that accounts for the low-level tracking error. Indeed, this constraint tightening strategy allows us to guarantee safety of the nonlinear system (2) in closed-loop with the proposed multi-rate control architecture, as we will discuss later on.

C. High Level Decision Making

In this section, we first describe how to compute a control policy which maximizes the probability of satisfying the specifications. Afterwards, we show how to compute the time-varying components p_{goal}^k , $\mathcal{C}_{\text{curr}}^k$, $\mathcal{C}_{\text{goal}}^k$, and $\mathcal{X}_{\text{goal}}^k$ used in the MPC problem (15).

Belief Model: For the MOMDP from Section III, we introduced the belief vector $b_k \in \mathcal{B}$ that represents the posterior probability that the partially observable state z_k equals $z \in \mathcal{Z}$.

Algorithm 1: Update High-Level

```

1 inputs:  $x(t)$ ,  $o_k$ ,  $s_{k-1}$ ,  $a_{k-1}$ ,  $b_{k-1}$  ;
2 set current high-level state  $s_k = \text{getState}(x(t))$  ;
3 compute current set  $\mathcal{C}_{\text{curr}}^k = \text{getCell}(s_k)$  ;
4 update belief  $b_k$  using (20) ;
5 compute high-level action  $a_k = \pi_h(s_k, b_k)$  ;
6 compute maximum likely estimate  $\hat{z}_k$  using (21);
7 update state  $s_{k+1} = f_s(s_k, \hat{z}_k, a_k)$  ;
8 compute goal set  $\mathcal{C}_{\text{goal}}^k = \text{getCell}(s_{k+1})$  ;
9 compute the forecasted action  $\hat{a} = \pi_h(s_{k+1}, b_k)$  ;
10 compute the forecasted state  $\hat{s}_{k+2} = f_s(s_{k+1}, \hat{z}_k, \hat{a})$  ;
11 set forecasted set  $\mathcal{C}_{\text{forc}}^k = \text{getCell}(s_{k+1})$  ;
12 get forecasted cell center  $c^{\text{forc}} = \text{getCenter}(\mathcal{C}_{\text{forc}}^k)$  ;
13 compute goal position  $p_{\text{goal}}^k = \text{Proj}(c^{\text{forc}}, \mathcal{C}_{\text{goal}}^k)$  ;
14 return:  $a_k$ ,  $b_k$ ,  $s_k$ ,  $\mathcal{C}_{\text{goal}}^k$ ,  $\mathcal{C}_{\text{curr}}^k$ ,  $p_{\text{goal}}^k$ 

```

The belief is a sufficient statistic and, for all $z' \in \mathcal{Z}$, it evolves accordingly to the following update equation:

$$b_{k+1}^{(z')} = \eta O(s_{k+1}, z', a_k, o_k) \times \sum_{z \in \mathcal{Z}} T_s(s_k, z, a_k, s_{k+1}) T_z(s_k, z, a_k, s_{k+1}, z_{k+1}) b_k^{(z)},$$

where η is a normalization constant [44]. Notice that the above update equation can be written in a compact form, i.e.,

$$b_{k+1} = f_b(s_{k+1}, s_k, o_k, a_k, b_k), \quad (20)$$

where $f_b : \mathcal{S} \times \mathcal{S} \times \mathcal{O} \times \mathcal{A} \times \mathcal{B} \rightarrow \mathcal{B}$. Finally, given the belief b_k , we introduce the following maximum likelihood environment state estimate:

$$\hat{z}_k = \underset{z \in \mathcal{Z}}{\text{argmax}} \mathbb{P}(z_k = z | o_k, s_k, a_{k-1}) = \underset{z \in \mathcal{Z}}{\text{argmax}} b^{(z)}. \quad (21)$$

Control Policy: At the highest level of abstraction our goal is to compute a control policy π_h , which maximizes the probability that the high-level trajectory ω satisfies the specifications ψ . Such control policy can be computed solving the following problem:

$$\pi_h = \underset{\pi}{\text{argmax}} \mathbb{P}^\pi[\omega \models \psi], \quad (22)$$

where $\mathbb{P}^\pi[\omega \models \psi]$ represents the probability that the specification ψ is satisfied for the closed-loop trajectory ω under the policy π . The solution to the above problem can be approximated using point-based and simulation-based strategies [6], [7], [17]–[19], [50]. In this work, we used the point-based strategy discussed in [42]. The resulting high-level control policy maps the high-level state s_k and the environment belief b_k to the high-level control action a_k , i.e., $a_k = \pi_h(s_k, b_k)$.

The high-level policy (22) is leveraged in Algorithm 1 to compute the goal position p_{goal}^k and the sets $\mathcal{C}_{\text{curr}}^k$ and $\mathcal{C}_{\text{goal}}^k$, which are used in the MPC problem (15). In Algorithm 1, we first use the function `getState`, which maps the current state $x(t)$ to the high-level state s_k representing the cell containing the nonlinear system 2 (line 2). Then, we compute the current cell $\mathcal{C}_{\text{curr}}^k$ associated with the high-level state s_k using the function `getCell` (line 3). Afterwards, we update the belief

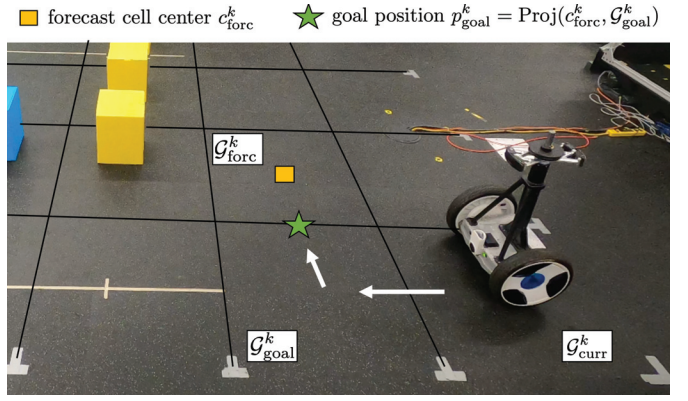


Fig. 3. The above figure illustrates the high-level updated from Algorithm 1 that is used to compute the goal position (green star).

state b_k and we compute the control action a_k (lines 4 – 5). Given the control action a_k and the maximum likelihood estimator of the environment state $\hat{z}_k$, we update the high-level state and we compute the goal cell $\mathcal{C}_{\text{goal}}^k$ (lines 6 – 8). Next, given the current belief b_k , we compute the action $\hat{a}$ that the high-level planner would select at the next update $k + 1$ assuming that the new observation is not informative, i.e., the belief $b_{k+1} = b_k$. We leverage the action $\hat{a}$ to estimate the high-level state $\hat{s}_{k+2}$, which represents the location where we should steer the system after transitioning to the high-level state s_{k+1} , if $b_{k+1} = b_k$. The state $\hat{s}_{k+2}$ is used to incorporate forecast into the high-level planner. In particular, given the $\hat{s}_{k+2}$ we compute the forecasted cell center $c^{\text{forc}} \in \mathbb{R}^{n_p}$ and the forecasted cell $\mathcal{C}_{\text{forc}}^k$ representing the grid cell where the high-level decision maker would like to steer the system, if no informative observations are collected (lines 11 – 12). Finally, the goal cell $\mathcal{C}_{\text{goal}}^k$ and the forecasted center $c^{\text{forc}} \in \mathbb{R}^{n_p}$ are used to compute the goal position p_{goal}^k (line 13).

Figure 3 illustrates the high-level update from Algorithm 1 that is used to compute the goal position leveraged in the design of the mid-level MPC. In this example, the Segway is located in the bottom right corner of the grid and the current high-level action a_k is to move west. The figure shows also the forecasted action $\hat{a}$ that the Segway would take from the goal region, if the belief b_k is not updated. Basically, $\hat{a}$ is a high-level open-loop prediction of the future control action and it is used to incorporate forecast into the high-level decision maker. Indeed, the goal position p_{goal}^k is computed projecting the forecasted cell center c^{forc} onto the goal cell $\mathcal{C}_{\text{goal}}^k$.

D. Control Architecture

Finally, we introduce the multi-rate hierarchical control architecture which leverages the low-level, mid-level, and high-level control policies from the previous sections. The multi-rate control Algorithm 2 describes the architecture depicted in Figure 2. When the nonlinear system (2) reaches the goal cell (i.e., $p(t) \in \mathcal{C}_{\text{goal}}^k$), the high-level decision maker reads the new observations o_{k+1} and updates high-level state, action, goal position p_{goal}^k , goal cell $\mathcal{C}_{\text{goal}}^k$, and current cell $\mathcal{C}_{\text{curr}}^k$ (lines 3–4). Finally, it updates the high-level time k and it initializes the MPC horizon $N_i^k = N$. Afterwards, the mid-level planner (lines 8 – 20) updates the mid-level time counter i and the

Algorithm 2: Multi-Rate Control

```

1 inputs:  $k, s_k, b_k, a_k, i, x(t), u_m(t), \bar{x}(t), C_{\text{curr}}^k, C_{\text{goal}}^k,$ 
    $p_{\text{goal}}^k, N_i^k, C_{\text{curr}}^{k-1}, C_{\text{goal}}^{k-1}, p_{\text{goal}}^{k-1}, N_i^{k-1}$  ;
2 if  $p(t) \in C_{\text{goal}}^k$  or  $k = 0$  then
   // Update high-level goal
3   measure  $o_{k+1}$  ;
4   update  $a_{k+1}, b_{k+1}, s_{k+1}, \mathcal{X}_{\text{goal}}^k, C_{\text{goal}}^{k+1}, C_{\text{curr}}^{k+1}, p_{\text{goal}}^{k+1}$ 
   using Algorithm 1 with  $x(t), o_{k+1}, s_k, a_k, b_k$  ;
5   set  $N_i^{k+1} = N$  ;
6    $k = k + 1$  ;
7 end
8 if  $t \in \mathcal{T}^c = \cup_{j=0}^{\infty} \{jT\}$  then
   // Update mid-level plan
9   solve MPC problem (15) with  $N = N_i^k$ , and  $\mathcal{X}_{\text{goal}}^k,$ 
    $C_{\text{curr}}^k, C_{\text{goal}}^k, p_{\text{goal}}^k$  ;
10  if the MPC problem (15) is not feasible then
11    solve MPC problem (15) with  $N = N_i^{k-1}$ , and
     $\mathcal{X}_{\text{goal}}^{k-1}, C_{\text{curr}}^{k-1}, C_{\text{goal}}^{k-1}, p_{\text{goal}}^{k-1}$  ;
12    set  $N_{i+1}^{k-1} = \max(1, N_i^{k-1} - 1)$  ;
13    set  $N_{i+1}^k = N_i^k$  ;
14  else
15    set  $N_{i+1}^{k-1} = N_i^{k-1}$  ;
16    set  $N_{i+1}^k = \max(1, N_i^k - 1)$  ;
17  end
18  set  $u_m(t) = v_{t|t}^{d,*} + K(x(t) - \bar{x}_{t|t}^{d,*})$  ;
19  update  $\bar{x}(t) = \Delta_{\bar{x}}(x(t)) = \bar{x}_{t|t}^{d,*}$  ;
20   $i = i + 1$  ;
21 end
   // Compute low-level control
22 solve the CBF problem (9) ;
23 Compute total input  $u(t) = u_l(t) + u_m(t)$  ;
24 Return:  $u(t), k, s_k, b_k, a_k, i, x(t), u_m(t), \bar{x}(t), C_{\text{curr}}^k,$ 
    $C_{\text{goal}}^k, p_{\text{goal}}^k, N_i^k, N_i^{k-1}$ 

```

planned trajectory at a constant frequency of $1/T$ Hz. First, it solves the MPC problem (15) with $N = N_i^k$ and time-varying components $\mathcal{X}_{\text{goal}}^k, C_{\text{goal}}^k, C_{\text{curr}}^k$, and p_{goal}^k . If the MPC problem is not feasible, the planner computes a contingency plan (lines 10 – 14), otherwise it updates the prediction horizon (lines 15-16). Note that the MPC problem solved in line 9 of Algorithm 1 may be not feasible as the terminal constraint set $\mathcal{X}_{\text{goal}}^k$ is updated by the high-level planner. For this reason, we introduced the contingency plan (lines 10–14), where the MPC problem from line 11 is constructed using the terminal constraint set $\mathcal{X}_{\text{goal}}^{k-1}$. As we will show in the proof of Theorem 1, when the MPC problem constructed with terminal constraint set $\mathcal{X}_{\text{goal}}^k$ is not feasible, we can guarantee the feasibility of the contingency MPC with $\mathcal{X}_{\text{goal}}^{k-1}$ as terminal constraint. This fact allows us to guarantee safety for the closed-loop system. Finally, Algorithm 2 computes the low-level control action solving the CLF-CBF QP (9) and the total control input that is given by the summation of the mid-level and low-level control actions, i.e.,

$$u(t) = u_l(t) + u_m(t).$$

V. SAFETY AND PERFORMANCE GUARANTEES

In this section we show the properties of the proposed multi-rate control architecture. We consider the augmented system:

$$\Sigma_{\text{aug}} : \begin{cases} \dot{x}(t) = f(x(t)) + g(x(t))(u_l(t) + u_m(t)), & t \geq 0 \\ \hat{x}(t) = A_{[t/T]} \bar{x}(t) + B_{[t/T]} u_m(t), & t \in \mathcal{T} \\ \bar{x}^+(t) = \Delta_{\bar{x}}(x(t)), & t \in \mathcal{T}^c \end{cases} \quad (23)$$

where the nonlinear dynamics for state $x(t) \in \mathbb{R}^{n_x}$ are defined in (2) and the LTV model for the nominal state $\bar{x}(t) \in \mathbb{R}^{n_x}$ is defined in (5) for the reset map (19) given by the MPC. In what follows, we analyse the properties of the proposed multi-rate control Algorithm 2 in closed-loop with system (23). We show that the closed-loop system satisfies state and input constraints (3) and that the proposed algorithm maximizes the probability of satisfying the specifications. Notice that in practice the state $x(t)$ is given by the nonlinear system (2), whereas the nominal state $\bar{x}(t)$ is computed by the low-level layer to update the tracking error $e(t)$, as shown in Figure 2.

Proposition 1. *Consider the closed-loop system (10) and (23) with mid-level input $u_m(t) \in \mathcal{U}_m$ and $\dot{u}_m(t) = 0, \forall t \in \mathcal{T}$. If Assumption 1 holds and the error $e(kT) = x(kT) - \bar{x}(kT) \in \mathcal{E}$ for all $k \in \mathbb{Z}_{0+}$, then the control policy (10) guarantees that $e(t) \in \mathcal{E}$ and $u_l(t) \in \mathcal{U}_l, \forall t \in [kT, (k+1)T)$.*

Proof: The proof follows from standard CBF arguments [28]. First, we notice that the error $e(kT) = x(kT) - \bar{x}(kT)$ follows the error dynamics in (6). Furthermore, by construction the time-varying matrices $(A_{[t/T]}, B_{[t/T]})$ are constant for $t \in [kT, (k+1)T)$. Therefore, for all $k \in \mathbb{Z}_{0+}$ and $t \in [kT, (k+1)T)$, we have that error dynamics in (6) are nonlinear control affine for the low-level input u_l . This fact implies that, if at time $t = kT$ the error $e(kT) = x(kT) - \bar{x}(kT) \in \mathcal{E}$, then from the feasibility of the CLF-CBF QP (9) from Assumption 1 we have that $e(t) = x(t) - \bar{x}(t) \in \mathcal{E}, \forall t \in [kT, (k+1)T)$. ■

Proposition 1 shows that between time $t_i = iT$ and $t_{i+1} = (i+1)T$ the difference between the state x and the state $\bar{x}$ of the reference trajectory is bounded. Next, we show that this property allows us to guarantee safety and convergence in finite time of the nonlinear system (2) to a goal cell C_{goal}^k contained in the feasible region. In turns, convergence in finite time allows us to show that the proposed approach maximizes the probability of satisfying the high-level specifications.

Assumption 3. Algorithm 1 returns a goal cell C_{goal}^k which is contained in the feasible set $\mathcal{X}_p$, i.e., $C_{\text{goal}}^k \subset \mathcal{X}_p$.

Theorem 1. *Let Assumptions 1-3 hold and consider system (23) in closed-loop with Algorithm 2. If at time $t_i = iT$ the MPC problem (15) is feasible with $N_i^k = N$ and time-varying components $\mathcal{X}_{\text{goal}}^k, C_{\text{curr}}^k, C_{\text{goal}}^k$, and p_{goal}^k , then there exists a $j \in \{i, \dots, i+N-1\}$ such that the closed-loop system satisfies state and input constraints (3) for all $t \in \{iT, \dots, jT\}$ and the state $x((j+1)T) = [p^{\top}((j+1)T), q^{\top}((j+1)T)]^{\top}$ reaches the goal cell C_{goal}^k , i.e., $p((j+1)T) \in C_{\text{goal}}^k$.*

Proof: From Assumption 3 we have that the high-level policy (22) takes a high-level action a_k which avoids collision

with the obstacles, i.e.,

$$\mathcal{C}_{\text{goal}}^k \subset \mathcal{X}_p. \quad (24)$$

Next, we show that if at time $t_i = iT$ the MPC problem (15) is feasible with $\mathcal{X}_{\text{goal}}^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k, p_{\text{goal}}^k$, and $N_i^k > 1$, then at time $t_{i+1} = (i+1)T$ the MPC problem (15) is feasible with $\mathcal{X}_{\text{goal}}^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k, p_{\text{goal}}^k$, and $N_{i+1}^k = N_i^k - 1$. Let

$$[x_{i|i}^{d,*}, x_{i+1|i}^{d,*}, \dots, x_{i+N_i^k|i}^{d,*}] \text{ and } [u_{i|i}^{d,*}, \dots, u_{i+N_i^k-1|i}^{d,*}] \quad (25)$$

be the optimal state input sequence to the MPC problem (15) at time $t_i = iT$. Then, from Proposition 1, equation (14), and the definition of the reset map (19), we have that

$$x((i+1)T) - \bar{x}_{i+1|i}^{d,*} = x((i+1)T) - \bar{x}^-((i+1)T) \in \mathcal{E}, \quad (26)$$

and therefore, by feasibility of (25) at time t_i , the following sequences of $N_i^k - 1$ states and $N_i^k - 2$ inputs

$$[x_{i+1|i}^{d,*}, \dots, x_{i+N_i^k|i}^{d,*}] \text{ and } [u_{i+1|i}^{d,*}, \dots, u_{i+N_i^k-1|i}^{d,*}] \quad (27)$$

are feasible at time $t_{i+1} = (i+1)T$ for the MPC problem (15) with $\mathcal{X}_{\text{goal}}^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k, p_{\text{goal}}^k$, and $N_{i+1}^k = N_i^k - 1$.

Now, we show that state and input constraints are satisfied until the system reaches the goal set $\mathcal{C}_{\text{goal}}^k$. Recall that by assumption the MPC problem is feasible at time $t_i = iT$ with $\mathcal{X}_{\text{goal}}^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k, p_{\text{goal}}^k, N_i = N$ and assume that $p(jT) \notin \mathcal{C}_{\text{goal}}^k$ for all $j \in \{i, \dots, i+N-1\}$. By induction the MPC problem (15) with $\mathcal{X}_{\text{goal}}^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k, p_{\text{goal}}^k$ and $N_j^k = N_i^k - j$ is feasible for all $j \in \{i, \dots, i+N-1\}$. Consequently, Algorithm 1 returns a feasible mid-level control action³ $u_m(t) \in \mathcal{U}_m$. Furthermore, from Proposition 1 we have that the low-level controller returns a feasible control action $u_l(t) \in \mathcal{U}_l$ and therefore

$$u(t) = u_l(t) + u_m(t) \in \mathcal{U}_l \oplus \mathcal{U}_m = \mathcal{U}, \forall t \in \mathbb{R}_{0+}. \quad (28)$$

The feasibility of the state-input sequences in (27) for the MPC problem solved with $\mathcal{X}_{\text{goal}}^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k, p_{\text{goal}}^k$ implies that

$$\begin{aligned} x_{j|j}^{d,*} &\in \mathcal{X}_{p,q}^k \ominus \mathcal{E} \\ x(jT) - x_{j|j}^{d,*} &\in \mathcal{E}, \end{aligned} \quad (29)$$

$\forall j \in \{i, \dots, i+N-1\}$. Consequently, from the above equation and definition (16), we have that

$$p(jT) \in \mathcal{X}_p \text{ and } q(jT) \in \mathcal{X}_q, \forall j \in \{i, \dots, i+N-1\}.$$

Finally, we show that the state $x(t)$ of the augmented system (23) in closed-loop with Algorithm 2 converges to the goal cell $\mathcal{C}_{\text{goal}}^k$ in finite time. We have shown that, if $p(jT) \notin \mathcal{C}_{\text{goal}}^k$ for all $j \in \{i, \dots, i+N-1\}$, then the MPC problem is feasible for all time $t_k = kT$ and $k \in \{i, \dots, i+N-1\}$. Now we notice that by feasibility of the MPC problem at time $t_{i+N-1} = (i+N-1)T$ with $N_{i+N-1} = 1$, we have that the optimal planned trajectory satisfies

$$x_{i+N|i+N-1}^{d,*} \in \mathcal{X}_{\text{goal}}^k \ominus \mathcal{E}_p, \quad (30)$$

where $\mathcal{E}_p$ is defined as in (17).

³Note that as $p(jT) \notin \mathcal{C}_{\text{goal}}^k$ for all $j \in \{i, \dots, i+N-1\}$ the MPC time-varying components are not updated.

From equation (14) and Proposition 1, we have that

$$x((i+N)T) - \bar{x}_{i+N|i+N-1}^{d,*} = x((i+N)T) - \bar{x}^-((i+N)T) \in \mathcal{E}.$$

The above equation together with (17) and (30) imply that at time $t_{i+N} = (i+N)T$

$$x((i+N)T) = \begin{bmatrix} p((i+N)T) \\ q((i+N)T) \end{bmatrix} \in \mathcal{X}_{\text{goal}}^k \ominus \mathcal{E}_p \oplus \mathcal{E}$$

and therefore $p((i+N)T) \in \mathcal{C}_{\text{goal}}^k$.

Concluding, if for all time $t_j = jT$ and $j \in \{i, \dots, i+N-1\}$ we have that $p(jT) \notin \mathcal{C}_{\text{goal}}^k$, then $p((i+N)T) \in \mathcal{C}_{\text{goal}}^k$. Thus, the closed-loop system converges to the goal cell $\mathcal{C}_{\text{goal}}^k$ in finite time. ■

Finally, we leverage Theorem 1 to show that, when $\mathcal{C}_{\text{goal}}^k \subset \mathcal{X}_p$, the multi-rate control Algorithm 2 steers the system in finite time to goal cell $\mathcal{C}_{\text{goal}}^k$ for all $k \in \mathbb{Z}_{0+}$ and, consequently, the closed-loop system maximizes the probability of satisfying the high-level specifications. In particular, we show that the contingency plan from lines 10–14 of Algorithm 2 guarantees feasibility of the planner when the time-varying components are updated.

Theorem 2. *Let Assumptions 1-3 hold and consider system (23) in closed-loop with Algorithm 2. If $x(0) \in \mathcal{X}_{\text{curr}}^k \oplus \mathcal{E}$, then Algorithm (2) maximizes the probability that the closed-loop satisfies the high-level specifications.*

Proof: The proof follows by induction. Assume that at time $t_i = iT$ the closed-loop system reaches the goal cell $\mathcal{C}_{\text{goal}}^k$, i.e., $p(iT) \in \mathcal{C}_{\text{goal}}^k$. Then, at time $t_i = iT$ we have that the high-level decision maker from Algorithm 2 (lines 2–9) updates the high-level time and the time-varying components $\mathcal{X}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{curr}}^{k+1}, \mathcal{C}_{\text{goal}}^{k+1}, p_{\text{goal}}^{k+1}$ used to design the MPC problem (15). After the high-level update, the MPC problem with $N = N_j^{k+1}, \mathcal{X}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{curr}}^{k+1}, \mathcal{C}_{\text{goal}}^{k+1}$, and p_{goal}^{k+1} may be either feasible or unfeasible⁴. Thus, we analyse the following three cases for $j \geq i$:

Case 1: The MPC problem with $\mathcal{C}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{curr}}^{k+1}, p_{\text{goal}}^{k+1}$, and $N = N_j^{k+1}$ is feasible, therefore from Theorem 1 we have that Algorithm 2 steers the nonlinear system to the goal $\mathcal{C}_{\text{goal}}^{k+1}$.

Case 2: The MPC problem with $\mathcal{C}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{curr}}^{k+1}, p_{\text{goal}}^{k+1}$, and $N = N_j^{k+1}$ is not feasible and $N_j^k = 1$. Then from Theorem 1, we have that the contingency MPC with $N_j^k, \mathcal{C}_{\text{goal}}^k, \mathcal{C}_{\text{curr}}^k$ and p_{goal}^k is feasible and Algorithm 1 returns a feasible control action. Furthermore, as $N_j^k = 1$ the terminal state of the optimal predicted trajectory is

$$x_{j+1|j}^{d,*} \in \mathcal{X}_{\text{goal}}^k \ominus \mathcal{E}_p.$$

The above equation together with equation (26) imply that

$$x((j+1)T) \in \mathcal{X}_{\text{goal}}^k \ominus \mathcal{E}_p \oplus \mathcal{E} \subset \mathcal{X}_{\text{goal}}^k \oplus \mathcal{E}.$$

Notice that $\mathcal{X}_{\text{goal}}^k = \mathcal{X}_{\text{curr}}^{k+1}$, thus from Assumption 2 we have that at the next time step $t_{j+1} = (j+1)T$ the MPC problem with $N_{j+1}^{k+1} = N, \mathcal{X}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{curr}}^{k+1}$, and p_{goal}^{k+1} is feasible and, from Theorem 1, we have that Algorithm 2 steers the nonlinear system to the goal $\mathcal{C}_{\text{goal}}^{k+1}$ in finite time.

⁴Unfeasibility may be caused by the update of $\mathcal{C}_{\text{goal}}^{k+1}, \mathcal{C}_{\text{curr}}^{k+1}$ and $\mathcal{X}_{\text{curr}}^{k+1}$.

Case 3: The MPC problem with $\mathcal{C}_{\text{goal}}^{k+1}$, $\mathcal{C}_{\text{curr}}^{k+1}$, p_{goal}^{k+1} , and $N = N_j^{k+1}$ is not feasible and $N_j^k > 1$. Then from Theorem 1, we have that the contingency MPC with N_j^k , $\mathcal{C}_{\text{goal}}^k$, $\mathcal{C}_{\text{curr}}^k$ and p_{goal}^k is feasible (lines 10–13 in Algorithm 2).

By assumption $x(0) \in \mathcal{X}_{\text{curr}}^k \oplus \mathcal{E}$, thus from Assumption 2 and Theorem 1 we have that at time $t = 0$ the MPC is feasible and that Algorithm 2 steers system (2) to $\mathcal{G}_{\text{goal}}^0$. This fact implies that the conditions form one the above Cases 1–3 are met, if $x(0) \in \mathcal{X}_{\text{curr}}^k \oplus \mathcal{E}$. Now notice that at each time step $N_{j+1}^k = N_j^k - 1$ (line 16), thus Case 3 occurs at most N times. Therefore, after at most N time steps the conditions from either Case 1 or Case 2 are met and Algorithm 2 will steer the system to the grid cell $\mathcal{C}_{\text{goal}}^{k+1}$ computed by the high-level policy (22). Consequently, as the high-level policy (22) maximizes the probability of satisfying the specifications, we have that the closed-loop system maximizes the probability of satisfying the specifications. ■

VI. RESULTS

We tested the proposed strategy in simulation and experiment on navigation tasks inspired by the Mars exploration mission [6], [7], [16]. We control a Segway-like robot and our goal is to explore the environment to find science samples which may be located in the known goal regions $\mathcal{G}_i$ shown in Figure 4. The specification $\psi = \neg \text{collision}U((\text{Goal}_1 \wedge \text{sample}_1) \vee (\text{Goal}_2 \wedge \text{sample}_2))$, where the atomic proposition sample_i is satisfied if the region $\mathcal{G}_i$ contains a science sample and the atomic proposition Goal_i is satisfied if the Segway is in a goal cell $\mathcal{G}_i$. The high-level control policy associated with specification ψ is computed solving a reach-avoid problem for the product MOMDP, which is computed preforming the cross-product between an automata associated

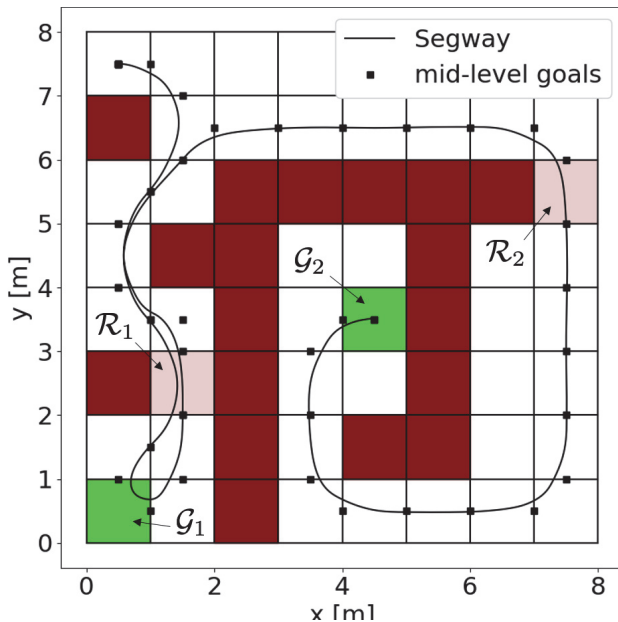


Fig. 4. Closed-loop trajectory. The Segway first explores regions $\mathcal{R}_1$, which is traversable, and $\mathcal{G}_1$ that does not contain the science sample. Afterwards, it explores the traversable region $\mathcal{R}_2$ and it reaches $\mathcal{G}_2$.

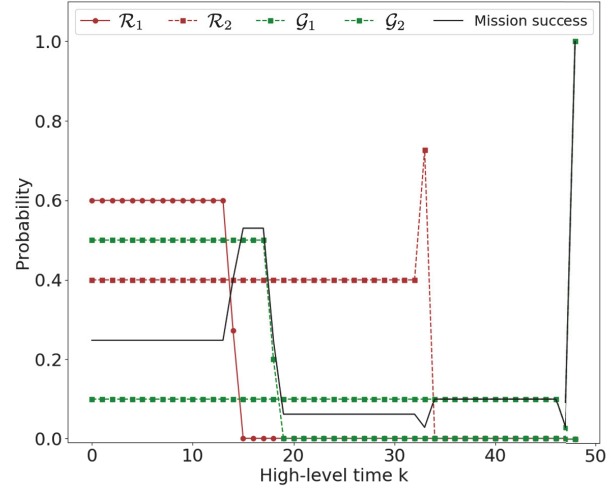


Fig. 5. This figure shows the closed-loop probability of mission success, which equals the probability of satisfying the high-level specifications. Furthermore, we reported also the belief about regions $\mathcal{R}_1$ and $\mathcal{R}_2$ being traversable (brown) and the goal regions $\mathcal{G}_1$ and $\mathcal{G}_2$ containing the science sample (green).

the specification ψ and the original MOMDP⁵. For further details on how to convert a specification into a finite state automata and the computation of the product, please refer to [43]. While performing the search task, we have to collect observations to determine the state of the uncertain region $\mathcal{R}_i$, which may be traversable with some probability. The controller has access only to partial observations about the environment. In particular, the Segway receives a perfect observation about the state of the uncertain region $\mathcal{R}_i$ when one cell away, an observation which is correct with probability 0.8, when the Manhattan distance is smaller than two, and an uninformative observations otherwise. Similarly, the Segway receives a partial noisy observation about the goal region $\mathcal{G}_i$ which is correct with probability 0.7, when one cell away and a perfect observations when the goal cell $\mathcal{G}_i$ is reached.

The state of the Segway is defined as follows:

$$x = [X, Y, \theta, v, \dot{\theta}, \psi, \dot{\psi}],$$

where (X, Y) represents the position of the center of mass, $(\theta, \dot{\theta})$ the heading angle and yaw rate, v the velocity, and $(\psi, \dot{\psi})$ the rod's angle and angular velocity. The control input $u = [T_l, T_r]$, where T_l and T_r are the torques to the left and right wheel motors, respectively. In order to implement the low-level CLF-CBF QP we used the following function:

$$h(e) = 1 - \|\text{diag}(v_h)(x - \bar{x})\|_2^2, \quad (31)$$

where $v_h = [1/0.02, 1/0.02, 1/0.1, 1/0.1, 1/0.3, 1/0.1, 1/0.3]$ and $\bar{x} = [\bar{X}, \bar{Y}, \bar{\theta}, \bar{v}, \bar{\theta}, \bar{\psi}, \bar{\psi}]$ represents the state of the

⁵The computational complexity of solving the high-level synthesis problem is a function of the dimension of the product MOMDP, which may grow exponentially for complex specifications. The analysis of the computational tractability of the high-level synthesis process is beyond the scope of this work and the code used to synthesize the high-level policy can be found at <https://github.com/urossolia/MOMDP>.

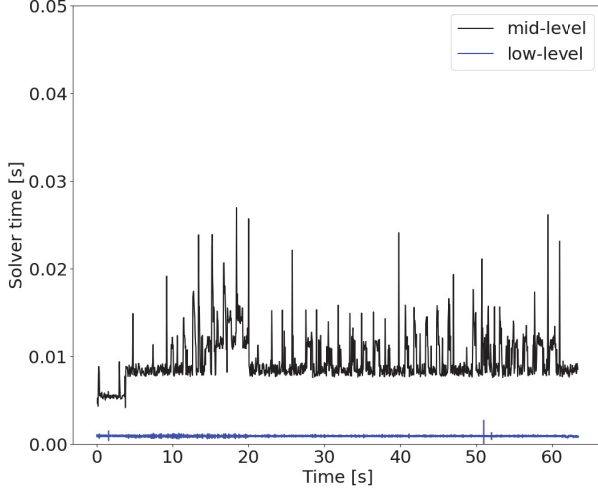


Fig. 6. Computational time associated with the mid and low layers. It takes on average 12 ms to compute mid-level control actions and less than 1 ms to compute low-level commands. In this example, the mid-level is discretized at 20 Hz and the low-level at 1 kHz.

nominal system from (5). The candidate control Lyapunov function is

$$V(e) = \|\text{diag}(v_v)(x - \bar{x})\|_2^2,$$

where $v_v = [100, 100, 100, 100, 10000, 10000, 100]$ and in the CLF-QBF QP (9) we used $c_1 = 1$, $c_2 = 10$ and $\alpha(x) = x$. The planning model (5) is computed iteratively linearizing the Segway dynamics around the predicted MPC trajectory. This strategy is standard in MPC, for more details on the linearization strategy please refer to [51]. The stage cost

$$\ell(x, u) = x^\top Q u + u^\top R u,$$

for the tuning matrices $Q = \text{diag}(0.1, 0.1, 0, 0, 10, 1, 10)$, $R = \text{diag}(0.01, 0.01)$, and $Q_f = \text{diag}(100, 100)$. Furthermore, we added an input rate cost with penalty $Q_{\text{rate}} = 0.1$ and a slack variable for the terminal constraint on the state $q_{t+N|i}$ with weight $Q_{\text{slack}} = \text{diag}(100, 100, 100)$. Finally, we approximated $\mathcal{S} = \{e = x - \bar{x} \in \mathbb{R}^{n_x} : h(e) \geq 0\} = \{e = x - \bar{x} \in \mathbb{R}^{n_x} : \|\text{diag}(v_v)(x - \bar{x})\|_2^2 \leq 1\}$ with $\tilde{\mathcal{S}} = \{e = x - \bar{x} \in \mathbb{R}^{n_x} : \|\text{diag}(v_v)(x - \bar{x})\|_\infty \leq 1\}$. This strategy allows us to write the MPC problem (15) as a QP⁶, which we solved using OSQP [52], [53].

A. Simulation

We implemented the proposed strategy in our high-fidelity Robotic Operating System (ROS) simulator. Figure 4 shows the locations of the uncertain and goal regions. The code can be found at https://github.com/DrewSingletary/segway_sim, please check the README.md to replicate our results. In this example the goal regions $\mathcal{G}_1$ and $\mathcal{G}_2$ may contain a science sample with probability 0.6 and 0.4, respectively. Whereas, regions $\mathcal{R}_1$ and $\mathcal{R}_2$ may be traversable with probability 0.5 and 0.1, as shown in Figure 5.

⁶Note that using $\mathcal{S}$ renders the MPC problem an SOCP, which is convex but computationally more demanding.

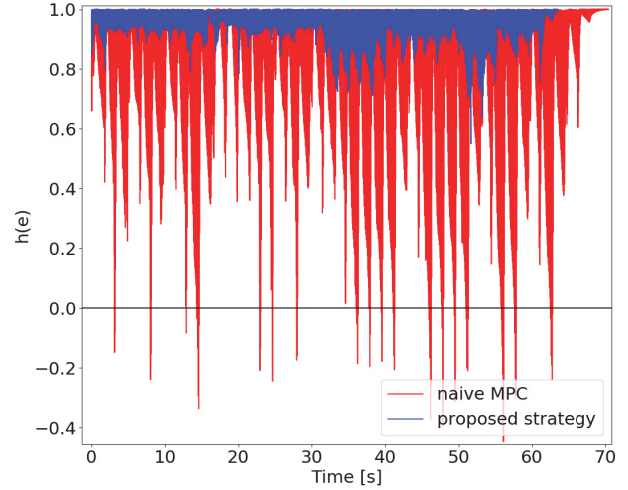


Fig. 7. Comparison between the barrier function associated with the proposed strategy and a naive MPC that is based on the linearized dynamics and it is synthesized without robustifying the constraints. As shown in the figure, when the low-level controller is not used the difference between the planner trajectory and the MPC trajectory grows and, as a results, the barrier function (31) becomes negative.

Figure 4 shows the closed-loop trajectory of the Segway. We notice that the controller explores the uncertain region $\mathcal{R}_1$, which in this example is traversable and afterwards it reaches the goal regions $\mathcal{G}_1$. As shown in Figure 5, at the high-level time $k = 19$ the controller figures out that the goal cell $\mathcal{G}_1$ does not contain a science sample and, consequently, the probability of mission success drops. Afterwards, the controller steers the Segway to the traversable region $\mathcal{R}_2$ and to the goal regions $\mathcal{G}_2$. In this example, the goal regions $\mathcal{G}_2$ contains a science sample and the mission is completed successfully, as shown in Figure 5.

The mid-level is discretized for $T = 50$ ms and the low-level at 1 kHz. Figure 6 shows the computational time associated with mid-level and low-level control actions. It takes on average 12 ms and at most about 30 ms to compute the mid-level control action $u_m(t)$ —thus the mid-level planner runs in real-time. Furthermore, we notice that it takes less than 1 ms to compute the low-level action $u_l(t)$.

Finally, we analyse the evolution of the barrier function (31), which quantifies the difference between the trajectory $x(t)$ of system (2) and the reference trajectory $\bar{x}(t)$ associated with nominal model (5), which is computed by the mid-level planner. We compared the proposed strategy with a naive MPC which is synthesized as in (15), but without taking into account the effect of the tracking error, i.e., we do not tighten the constraints and we set $x_{i|i} = x(t)$. Figure 7 shows the evolution of the barrier function for the proposed strategy and the naive MPC. We notice that when the low-level controller is not used, the barrier function becomes negative and in general has a lower magnitude. Therefore, this figure shows the advantage of the proposed hierarchical control architecture, where the low-level high-frequency controller is leveraged to track the reference trajectory. Indeed, this high-frequency

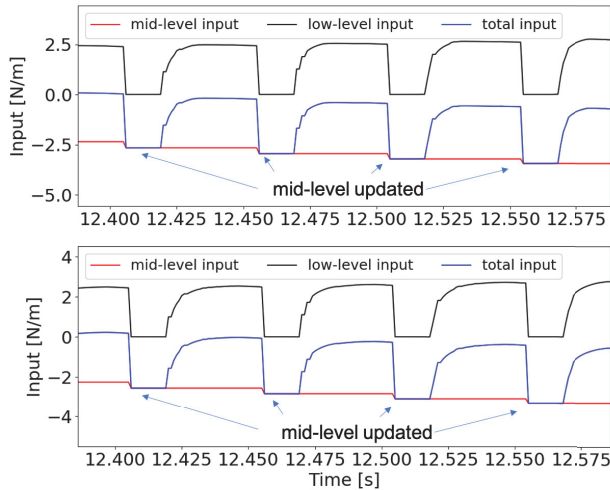


Fig. 8. Input torque sent to the right (top) and left (bottom) motor over a period of 0.2 second. The mid-level input is updated at 20 Hz, whereas the low-level action is updated at 1 kHz.

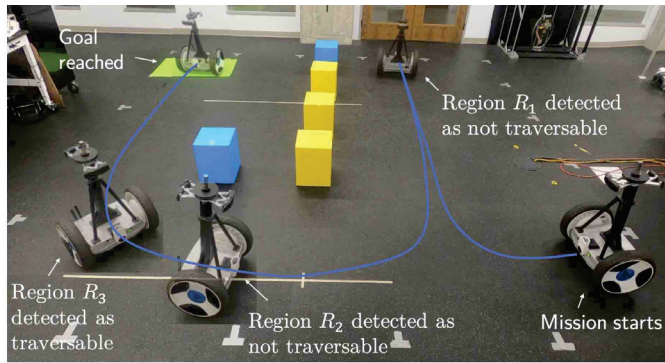


Fig. 9. Closed-loop trajectory during the experiment. The Segway first explores the uncertain regions $\mathcal{R}_1$, $\mathcal{R}_2$, and $\mathcal{R}_3$ and afterwards it reaches the goal region.

feedback is used to modify the mid-level control actions, as shown in Figure 8. The mid-level control action is updated at 20 Hz and the low-level input at 1 kHz. Notice that after the update of the mid-level input, the contribution of the low-level input towards the total control action is limited. However, as time progresses the accuracy of the linearization used to plan the reference trajectory decreases and the effort required by low-level controller to track the trajectory computed by the mid-level planner increases.

B. Experiment

We implemented the proposed multi-rate hierarchical control strategy on the Segway-like robot shown in Figure 1. State estimation is based on wheel encoders and IMU data from a VectorNav VN-100. The state estimate and the low-level control action u_l are computed at 800 Hz on the Segway, which is equipped with an ARM Cortex-A57 (quad-core) @ 2 GHz CPU running the ERIKA3 RTOS. On the other hand, the mid-level planner discretized at 20 Hz and the high-level decision maker run on a desktop with an Intel Core i7-8700

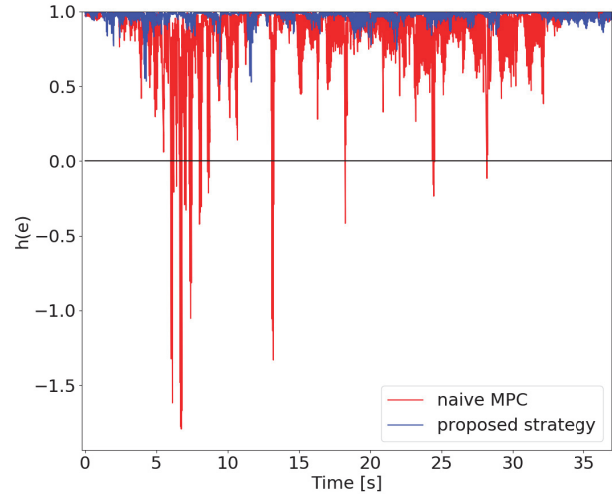


Fig. 10. Experimental comparison between the barrier function associated with the proposed strategy and a naive MPC which is based on the linearized dynamics. Also in this case, when the low-level controller is not used, the difference between the planner trajectory and the MPC trajectory grows and, as a results, the barrier function defined in (31) becomes negative.

CPU (6-cores) @ 3.7 GHz CPU, which sends the reference trajectory $\bar{x}$ and the reference input u_m via WiFi.

Figure 1 shows the location of the three uncertain regions $\mathcal{R}_1$, $\mathcal{R}_2$ and $\mathcal{R}_3$ which may be traversable with probability 0.9, 0.3 and 0.2, respectively. In this example, we assume that the goal region $\mathcal{G}_1$ contains the science sample with probability 1. Figure 9 shows the closed-loop trajectory. First, the controller explores region $\mathcal{R}_1$ which is not traversable, and afterwards it steers the Segway towards regions $\mathcal{R}_2$ and $\mathcal{R}_3$. After collecting observations about the environment, the controller detects that region $\mathcal{R}_2$ is not traversable and that region $\mathcal{R}_3$ is free space that the Segway can navigate through to reach the goal region $\mathcal{G}_1$. A video of the experiment and comparison with a naive MPC can be found at https://www.youtube.com/watch?v=Q-Mm0YwPh_I.

Figure 10 shows the evolution of the control barrier function (31). We compare the proposed strategy with a naive MPC which is designed as in (15), but without robustifying the constraint sets and setting $x_{i|i} = x(t)$. Also in this case, when the high-frequency low-level controller is not active, the barrier function becomes negative meaning that the error e does not belong to the safe set $\mathcal{E}$, i.e., $e(t) \notin \mathcal{E}$ for all $t \in \mathbb{R}_{0+}$. This result highlights the importance of the low-level high-frequency feedback from the CLF-CBF QP, which compensates for the model mismatch at the planning layer. Indeed, the MPC planner uses a linearized and discretized model, which is a first order approximation of the true dynamics. This approximation is accurate only at the discrete time instances when the MPC input is computed. To compensate for this model inaccuracy, the low-level CLF-CBF QP tracking controller computes the high-frequency component $u_l(t)$ that is added to the mid-level piecewise constant input $u_m(t)$, as shown in Figure 11.

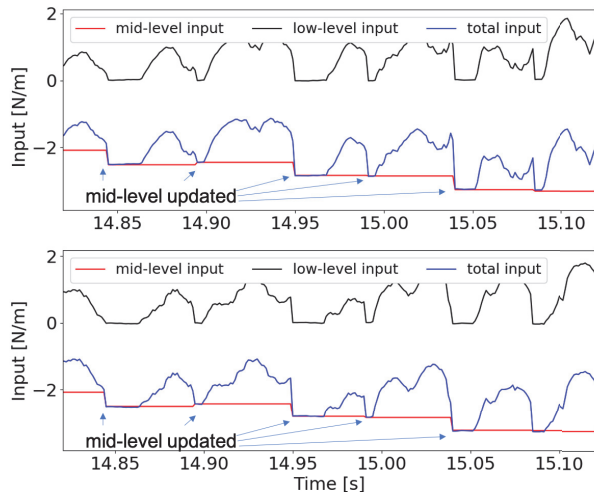


Fig. 11. Experimental results. Input torque sent to the right (top) and left (bottom) motor over a period of 0.3 seconds. The mid-level input is updated at 20 Hz, whereas the low-level action is updated at 800 Hz. Notice that the total input is the summation of the low-level and mid-level control inputs.

VII. CONCLUSIONS

In this paper we presented a multi-rate hierarchical control architecture for navigation tasks in partially observable environments. At the lowest level we leverage a CLF-CBF QP, which is used to track a reference trajectory within some error bounds. The reference trajectory is computed by a mid-level planner which leverages an MPC with time-varying terminal components. The feasibility of the MPC planner is guaranteed via a contingency scheme and a local reachability assumption on the planning model. Finally, at the highest level of abstraction, we showed how to model the system-environment interaction using a MOMDP and we proposed an algorithm to update the MPC time-varying components. The effectiveness of the proposed strategy is shown on navigation examples, where a Segway-like robot has to find science samples, while avoiding partially observable obstacles.

VIII. ACKNOWLEDGEMENTS

The authors would like to thank Geoffroy le Courtois du Manoir for helping with experiments and anonymous reviewers for constructive suggestions.

REFERENCES

- [1] T. Wongpiromsarn, U. Topcu, N. Ozay, H. Xu, and R. M. Murray, "Tulip: a software toolbox for receding horizon temporal logic planning," in *Proceedings of the 14th international conference on Hybrid systems: computation and control*, 2011, pp. 313–314.
- [2] T. Wongpiromsarn, U. Topcu, and R. M. Murray, "Receding horizon temporal logic planning," *IEEE Transactions on Automatic Control*, vol. 57, no. 11, pp. 2817–2830, 2012.
- [3] T. Wongpiromsarn, U. Topcu, and R. M. Murray, "Receding horizon control for temporal logic specifications," in *Proceedings of the 13th ACM international conference on Hybrid systems: computation and control*, 2010, pp. 101–110.
- [4] P. Tabuada and G. J. Pappas, "Linear time logic control of discrete-time linear systems," *IEEE Transactions on Automatic Control*, vol. 51, no. 12, pp. 1862–1877, 2006.

- [5] R. Alur, T. A. Henzinger, G. Lafferriere, and G. J. Pappas, "Discrete abstractions of hybrid systems," *Proceedings of the IEEE*, vol. 88, no. 7, pp. 971–984, 2000.
- [6] S. Haesaert, R. Thakker, R. Nilsson, A. Agha-mohammadi, and R. M. Murray, "Temporal logic planning in uncertain environments with probabilistic roadmaps and belief spaces," in *2019 IEEE 58th Conference on Decision and Control (CDC)*. IEEE, 2019, pp. 6282–6287.
- [7] S. Haesaert, P. Nilsson, C. I. Vasile, R. Thakker, A.-a. Agha-mohammadi, A. D. Ames, and R. M. Murray, "Temporal logic control of POMDPs via label-based stochastic simulation relations," *IFAC-PapersOnLine*, vol. 51, no. 16, pp. 271–276, 2018.
- [8] S. Kousik, S. Vaskov, F. Bu, M. Johnson-Roberson, and R. Vasudevan, "Bridging the gap between safety and real-time performance in receding-horizon trajectory design for mobile robots," *arXiv preprint arXiv:1809.06746*, 2018.
- [9] Y. S. Shao, C. Chao, S. Kousik, and R. Vasudevan, "Reachability-based trajectory safeguard (rts): A safe and fast reinforcement learning safety layer for continuous control," *arXiv preprint arXiv:2011.08421*, 2020.
- [10] S. L. Herbert, M. Chen, S. Han, S. Bansal, J. F. Fisac, and C. J. Tomlin, "FaSTrack: A modular framework for fast and guaranteed safe motion planning," in *2017 IEEE 56th Annual Conference on Decision and Control (CDC)*. IEEE, 2017, pp. 1517–1522.
- [11] N. Csomay-Shanklin, A. J. Taylor, U. Rosolia, and A. D. Ames, "Multi-rate planning and control of uncertain nonlinear systems: Model predictive control and control lyapunov functions," *arXiv preprint arXiv:2204.00152*, 2022.
- [12] A. Pnueli, "The temporal logic of programs," in *18th Annual Symposium on Foundations of Computer Science (sfcs 1977)*. IEEE, 1977, pp. 46–57.
- [13] S. G. Loizou and K. J. Kyriakopoulos, "Automatic synthesis of multi-agent motion tasks based on LTL specifications," in *2004 43rd IEEE Conference on Decision and Control (CDC)(IEEE Cat. No. 04CH37601)*, vol. 1. IEEE, 2004, pp. 153–158.
- [14] G. E. Fainekos, H. Kress-Gazit, and G. J. Pappas, "Hybrid controllers for path planning: A temporal logic approach," in *Proceedings of the 44th IEEE Conference on Decision and Control*. IEEE, 2005, pp. 4885–4890.
- [15] M. Kloetzer and C. Belta, "A fully automated framework for control of linear systems from temporal logic specifications," *IEEE Transactions on Automatic Control*, vol. 53, no. 1, pp. 287–297, 2008.
- [16] P. Nilsson, S. Haesaert, R. Thakker, K. Otsu, C.-I. Vasile, A.-A. Agha-Mohammadi, R. M. Murray, and A. D. Ames, "Toward specification-guided active mars exploration for cooperative robot teams," *Robotics: Science and Systems (RSS)*, 2018.
- [17] M. Bouton, J. Tumova, and M. J. Kochenderfer, "Point-based methods for model checking in partially observable markov decision processes," in *AAAI*, 2020, pp. 10 061–10 068.
- [18] C.-I. Vasile, K. Leahy, E. Cristofalo, A. Jones, M. Schwager, and C. Belta, "Control in belief space with temporal logic specifications," in *2016 IEEE 55th Conference on Decision and Control (CDC)*. IEEE, 2016, pp. 7419–7424.
- [19] Y. Wang, S. Chaudhuri, and L. E. Kavraki, "Bounded policy synthesis for POMDPs with safe-reachability objectives," *arXiv preprint arXiv:1801.09780*, 2018.
- [20] M. Ahmadi, R. Sharan, and J. W. Burdick, "Stochastic finite state control of POMDPs with LTL specifications," *arXiv preprint arXiv:2001.07679*, 2020.
- [21] M. Kwiatkowska, G. Norman, and D. Parker, "PRISM 4.0: Verification of probabilistic real-time systems," in *International conference on computer aided verification*. Springer, 2011, pp. 585–591.
- [22] C. Dehnert, S. Junges, J.-P. Katoen, and M. Volk, "A storm is coming: A modern probabilistic model checker," in *International Conference on Computer Aided Verification*. Springer, 2017, pp. 592–600.
- [23] E. Altman, *Constrained Markov decision processes*. CRC Press, 1999, vol. 7.
- [24] E. J. Sondik, "The optimal control of partially observable markov processes over the infinite horizon: Discounted costs," *Operations research*, vol. 26, no. 2, pp. 282–304, 1978.
- [25] P. Poupart and C. Boutilier, "Bounded finite state controllers," in *NIPS*, 2003.
- [26] J. Pineau, G. Gordon, S. Thrun *et al.*, "Point-based value iteration: An anytime algorithm for POMDPs," in *IJCAI*, vol. 3, 2003, pp. 1025–1032.
- [27] T. Gurriet, A. Singletary, J. Reher, L. Ciarletta, E. Feron, and A. Ames, "Towards a framework for realizable safety critical control through active set invariance," in *2018 ACM/IEEE 9th International Conference on Cyber-Physical Systems (ICCPs)*. IEEE, 2018, pp. 98–106.

- [28] 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*, vol. 62, no. 8, pp. 3861–3876, Aug 2017.
- [29] L. Wang, A. D. Ames, and M. Egerstedt, "Safety barrier certificates for collisions-free multirobot systems," *IEEE Transactions on Robotics*, vol. 33, no. 3, pp. 661–674, 2017.
- [30] K. P. Wabersich and M. N. Zeilinger, "Linear model predictive safety certification for learning-based control," in *2018 IEEE Conference on Decision and Control (CDC)*. IEEE, 2018, pp. 7130–7135.
- [31] H. Yin, M. Bujarbaruah, M. Arcak, and A. Packard, "Optimization based planner tracker design for safety guarantees," *arXiv preprint arXiv:1910.00782*, 2019.
- [32] S. Singh, M. Chen, S. L. Herbert, C. J. Tomlin, and M. Pavone, "Robust tracking with model mismatch for fast and safe planning: an sos optimization approach," in *International Workshop on the Algorithmic Foundations of Robotics*. Springer, 2018, pp. 545–564.
- [33] S. Singh, A. Majumdar, J.-J. Slotine, and M. Pavone, "Robust online motion planning via contraction theory and convex optimization," in *2017 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, 2017, pp. 5883–5890.
- [34] Y. Gao, A. Gray, H. E. Tseng, and F. Borrelli, "A tube-based robust nonlinear predictive control approach to semiautonomous ground vehicles," *Vehicle System Dynamics*, vol. 52, no. 6, pp. 802–823, 2014.
- [35] M. Kögel and R. Findeisen, "Discrete-time robust model predictive control for continuous-time nonlinear systems," in *2015 American Control Conference (ACC)*. IEEE, 2015, pp. 924–930.
- [36] S. Yu, C. Maier, H. Chen, and F. Allgöwer, "Tube MPC scheme based on robust control invariant set with application to lipschitz nonlinear systems," *Systems & Control Letters*, vol. 62, no. 2, pp. 194–200, 2013.
- [37] D. Gundana and H. Kress-Gazit, "Event-based signal temporal logic synthesis for single and multi-robot tasks," *IEEE Robotics and Automation Letters*, vol. 6, no. 2, pp. 3687–3694, 2021.
- [38] J. Chen, H. Wang, M. Rubenstein, and H. Kress-Gazit, "Automatic control synthesis for swarm robots from formation and location-based high-level specifications," in *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2020, pp. 8027–8034.
- [39] J. Köhler, R. Soloperto, M. A. Muller, and F. Allgower, "A computationally efficient robust model predictive control framework for uncertain nonlinear systems," *IEEE Transactions on Automatic Control*, 2020.
- [40] U. Rosolia and A. D. Ames, "Multi-rate control design leveraging control barrier functions and model predictive control policies," *IEEE Control Systems Letters*, vol. 5, no. 3, pp. 1007–1012, 2021.
- [41] D. Q. Mayne, M. M. Seron, and S. Raković, "Robust model predictive control of constrained linear systems with bounded disturbances," *Automatica*, vol. 41, no. 2, pp. 219–224, 2005.
- [42] U. Rosolia, M. Ahmadi, R. M. Murray, and A. D. Ames, "Time-optimal navigation in uncertain environments with high-level specifications," *To appear on the IEEE Conference on Decision and Control (CDC)*, *arXiv preprint arXiv:2103.01476*, 2021.
- [43] C. Belta, B. Yordanov, and E. A. Gol, *Formal methods for discrete-time dynamical systems*. Springer, 2017, vol. 89.
- [44] S. C. Ong, S. W. Png, D. Hsu, and W. S. Lee, "Planning under uncertainty for robotic tasks with mixed observability," *The International Journal of Robotics Research*, vol. 29, no. 8, pp. 1053–1068, 2010.
- [45] Y. Chen, H. Peng, J. Grizzle, and N. Ozay, "Data-driven computation of minimal robust control invariant set," in *2018 IEEE Conference on Decision and Control (CDC)*. IEEE, 2018, pp. 4052–4058.
- [46] S. Karaman and E. Frazzoli, "Sampling-based algorithms for optimal motion planning," *The international journal of robotics research*, vol. 30, no. 7, pp. 846–894, 2011.
- [47] O. Arslan and P. Tsiotras, "Use of relaxation methods in sampling-based algorithms for optimal motion planning," in *2013 IEEE International Conference on Robotics and Automation*. IEEE, 2013, pp. 2421–2428.
- [48] D. Ghosh, G. Nandakumar, K. Narayanan, V. Honkote, and S. Sharma, "Kinematic constraints based Bi-directional RRT (KB-RRT) with parameterized trajectories for robot path planning in cluttered environment," in *2019 International Conference on Robotics and Automation (ICRA)*. IEEE, 2019, pp. 8627–8633.
- [49] J. J. Kuffner and S. M. LaValle, "Rrt-connect: An efficient approach to single-query path planning," in *Proceedings 2000 ICRA. Millennium Conference. IEEE International Conference on Robotics and Automation. Symposia Proceedings (Cat. No. 00CH37065)*, vol. 2. IEEE, 2000, pp. 995–1001.
- [50] G. Shani, J. Pineau, and R. Kaplow, "A survey of point-based pomdp solvers," *Autonomous Agents and Multi-Agent Systems*, vol. 27, no. 1, pp. 1–51, 2013.

- [51] U. Rosolia and F. Borrelli, "Learning how to autonomously race a car: a predictive control approach," *IEEE Transactions on Control Systems Technology*, 2019.
- [52] B. Stellato, G. Banjac, P. Goulart, A. Bemporad, and S. Boyd, "Osqp: An operator splitting solver for quadratic programs," *Mathematical Programming Computation*, pp. 1–36, 2020.
- [53] G. Banjac, B. Stellato, N. Moehle, P. Goulart, A. Bemporad, and S. Boyd, "Embedded code generation using the osqp solver," in *2017 IEEE 56th Annual Conference on Decision and Control (CDC)*. IEEE, 2017, pp. 1906–1911.



Ugo Rosolia received the B.S. and M.S. (*cum laude*) degrees in mechanical engineering from the Politecnico di Milano, Milan, Italy, in 2012 and 2014, and a Ph.D. from the University of California Berkeley in 2019. He is currently a postdoctoral scholar with the California Institute of Technology.

He was a Visiting Scholar with Tongji University, Shanghai, China, for the Double Degree Program PoliTong from 2010 to 2011. During his M.S. degree, he was a Visiting Student for two semesters with the University of Illinois at Urbana-Champaign, Urbana, IL, USA, sponsored by a Global E3 Scholarship. He was a Research Engineer with Siemens PLM Software, Leuven, Belgium, in 2015, where he was involved in the optimal control algorithms. His current research interests include nonlinear optimal control for the centralized and decentralized system, the iterative learning control, and the predictive control.



Andrew Singletary received his B.S. degrees in Mechanical Engineering and Nuclear and Radiological Engineering from the Georgia Institute of Technology in 2017. He is currently at the California Institute of Technology, where he received his M.S. in Mechanical Engineering in 2019, and where he is currently pursuing a Ph.D. in Mechanical Engineering. His current research interests include safety critical control and planning for multi-agent systems.



Aaron D. Ames (S'04–M'16) received a B.S. in mechanical engineering and a B.A. in mathematics from the University of St. Thomas in 2001, and he received a M.A. in mathematics and a Ph.D. in electrical engineering and computer sciences from UC Berkeley in 2006. He is the Bren Professor of Mechanical and Civil Engineering and Control and Dynamical Systems at Caltech. Prior to joining Caltech in 2017, he was an Associate Professor at Georgia Tech in the Woodruff School of Mechanical Engineering

and the School of Electrical & Computer Engineering. He served as a Postdoctoral Scholar in Control and Dynamical Systems at Caltech from 2006 to 2008, and began his faculty career at Texas A&M University in 2008. His research interests span the areas of robotics, nonlinear control and hybrid systems, with a special focus on applications to bipedal robotic walking both formally and through experimental validation. His lab designs, builds and tests novel bipedal robots, humanoids and prostheses with the goal of achieving human-like bipedal robotic locomotion and translating these capabilities to robotic assistive devices. Dr. Ames received the 2005 Leon O. Chua Award for achievement in nonlinear science, the 2006 Bernard Friedman Memorial Prize in Applied Mathematics, the NSF CAREER award in 2010, and the 2015 Donald P. Eckman Award.