

A Formal Safety Net for Waypoint Following in Ground Robots

Brandon Bohrer^{1,2}, Yong Kiam Tan¹, Stefan Mitsch¹, Andrew Sogokon¹, and André Platzer^{1,2}

Abstract—We present a reusable formally verified safety net that provides end-to-end safety and liveness guarantees for 2D waypoint-following of Dubins-type ground robots with tolerances and acceleration. We: *i*) Model a robot in differential dynamic logic (dL), and specify assumptions on the controller and robot kinematics, *ii*) Prove formal safety and liveness properties for waypoint-following with speed limits, *iii*) Synthesize a monitor, which is automatically proven to enforce model compliance at runtime, and *iv*) Our use of the VeriPhy toolchain makes these guarantees carry over down to the level of machine code with untrusted controllers, environments, and plans. The guarantees for the safety net apply to any robot as long as the waypoints are chosen safely and the physical assumptions in its model hold. Experiments show these assumptions hold in practice, with an inherent trade-off between compliance and performance.

Index Terms—Formal Methods in Robotics and Automation, Robot Safety, Hybrid Logical/Dynamical Planning and Verification, Motion Control, Kinematics

I. INTRODUCTION

MANY autonomous ground robots are safety-critical because they operate near or in concert with humans. Formally verifying these systems is important: logic allows rigorous correctness arguments that apply in *all* system states, providing a powerful complement to system testing. Yet for robotics, even choosing a property to formally verify is challenging: many modeling abstractions and safety properties are available, with competing trade-offs. Discrete techniques can be applied to control software, but robots are *cyber-physical systems*: their verification must account for discrete controllers, continuous mechanics, and interactions between them.

Manuscript received: January, 21, 2019; Revised April, 23, 2019; Accepted May, 29, 2019.

This paper was recommended for publication by Editor Dezhen Song upon evaluation of the Associate Editor and Reviewers' comments. This work was supported by NDSEG, A*STAR, AFOSR grant FA9550-16-1-0288, NSF CPS Award CNS-1739629, the United States Air Force and DARPA under Contract No. FA8750-18-C-0092, and the Alexander von Humboldt Foundation. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of any sponsoring institution, the U.S. government or any other entity.

¹The authors are with Computer Science Department, Carnegie Mellon University, Pittsburgh, USA. {bbohrer, yongkiat, smitsch, asogokon, aplatzer}@cs.cmu.edu

²Part of this work was performed while the first and last authors were at Fakultät für Informatik, Technische Universität München.

Digital Object Identifier (DOI): 10.1109/LRA.2019.2923099.

©2019 IEEE. Personal use of this material is permitted. Permission from IEEE must be obtained for all other uses, in any current or future media, including reprinting/republishing this material for advertising or promotional purposes, creating new collective works, for resale or redistribution to servers or lists, or reuse of any copyrighted component of this work in other works.

Hybrid systems [1] emerged as mathematical models for robots because they integrate discrete with continuous dynamics, but they raise new questions about what it means to verify a robot. Robot kinematics are endlessly nuanced, so any model is always an approximation of reality. Even control software is rarely modeled exactly: simplifications are necessary in practice to limit verification complexity. Moreover, control software evolves throughout its development, and since verification of arbitrary programs cannot be fully automated, re-verifying control code after every change would be impractical.

How then ought a robot be verified? Online monitoring, per the Simplex method [2], offers a solution: run the control software, but treat its control decision as an *untrusted suggestion* which is supervised against a trusted *monitor condition* ϕ describing all “known-safe” decisions. Whenever the suggestion is known-safe the supervising *monitor* allows it to proceed, but otherwise invokes a trusted fallback decision, like emergency braking, to regain safety before returning control to the untrusted controller. Online monitoring is appealing because it enables treating the controller as a black box: only the monitor condition and fallback are safety-critical, both of which are simpler and can often be re-used as the control software evolves, or even across hardware platforms. To promote reusability, we target a waypoint-following notion of safety: other notions of safety such as collision-avoidance often come with restrictive assumptions, e.g., on the quantity and dynamics of obstacles, whereas waypoint-following abstracts away such problems under the choice of safe waypoints.

The monitor conditions and fallback are both best kept simple. While steering enables active safety and may reduce the required braking power, we simply brake at the maximum rate. Regardless which fallback is used, however, what is essential is that the monitor conditions and fallback provide safety. The ModelPlex [3] synthesizer and VeriPhy [4] compilation toolchain ensures safety at implementation level by: *i*) synthesizing correct-by-construction *monitor conditions* ϕ from a proven-safe hybrid systems model containing a proven-safe fallback (ModelPlex) *ii*) soundly compiling high-level monitor conditions and high-level fallback programs to machine-code monitor implementations with sound machine arithmetic (VeriPhy). ModelPlex and VeriPhy expand upon the reusability inherent to the black-box approach: hybrid systems models and proofs can treat many system parameters (e.g., tolerances and system delay) *generically*, once and for all, for all choices of the parameters. The hybrid system model can be used as a template: new monitors can often be generated for new systems *without* doing new proofs, so long as choosing

system parameters suffices to faithfully model the new system. The VeriPhy approach is also *end-to-end* in that VeriPhy outputs a chain of formal proofs, in theorem provers, that the actions taken by the machine-level (monitored) program fall under the original model and are thus formally safe. Even then, some guarantees are lost when sensors and actuators are buggy, when monitor conditions for the physical dynamics are violated, or when an unsafe plan is provided: we discuss these limitations in related work (Sect. II).

The VeriPhy approach has until now only been tested on an overly simplistic low-speed robot implementing a model of straight-line motion with direct velocity control. In this paper, we show for the first time that this approach scales to realistic models and simulations:

- We model as a hybrid program [5] a 2D accelerating vehicle that follows bloated Dubins paths.
- We prove safety and liveness properties for the model with the KeYmaera X theorem prover [6] for differential dynamic logic (dL) [5].
- We use VeriPhy [4] to synthesize a monitor and to automatically prove its correctness down to the machine code implementation, which gives an end-to-end proof.

Waypoint-following has the advantage of a clean interface to other robot software with its dual purposes of both *safety* (avoiding unsafe regions) and *liveness* (reaching its goal). We call waypoint-following safe iff the robot always follows the given path to its waypoint within a given tolerance and obeys given speed limits. Collision freedom then reduces to checking that correct waypoints and speed limits were given. The system is *live* if at all points it is possible to drive the rest of the way to the waypoint.

Obstacle avoidance [7], in contrast, directly verifies collision freedom, but liveness is challenging to even state let alone prove. The mix of safety and liveness is essential because a motionless robot is technically safe, but neither live nor useful because it never reaches its goal. By studying safety and liveness of waypoint-following, we provide a clean separation of concerns compared to the orthogonal question of verifying a discrete planner [8].

Bloating the 2D Dubins dynamics adds an additional tolerance margin to the ideal dynamics which accounts for the gap between the approximate dynamical model and reality (Sect. IV). Our evaluation shows that: *i*) a variety of classical control choices such as bang-bang and PD control fit within the bloated ideal path *ii*) the model assumptions hold in practice because AirSim's non-holonomic dynamics fall within bloated ideal holonomic Dubins dynamics, and *iii*) there is a trade-off between meeting model assumptions and operational performance: more aggressive controllers break assumptions more often. This paper also serves as a case study on safety and liveness verification: once the *safety* property is proved, much of the effort can be reused to prove *liveness*. The safety and liveness proof were performed interactively, but crucially need only be performed once per dL model, which describes an entire class of systems. Thanks to the automated proofs provided by VeriPhy, runtime monitors can be applied to new controls and even new hardware or simulation platforms with no additional manual proofs, so long as the controls and

dynamics stay within a tolerance around the ideal holonomic dynamics, so that the same dL model applies. Because the model treats system parameters (such as system delays and tolerances) generically, a wide variety of Dubins-like systems are already supported simply by changing the parameters. That is how we developed a formal safety net for ground robots and evaluated it on a realistic simulation. Because the monitor is reusable, we hope our safety net can assist future implementers in developing new systems.

II. RELATED WORK

Related work in formal methods and robotics applied synthesis and verification techniques to safe robotic control. This paper is the first to use a verified-safe monitor to enforce waypoint-following correctness of a realistic simulation.

A. Synthesis for Verified Planning and Control

Much of the existing related work considers high-level plan synthesis in isolation, with informal proofs of correctness. Our work is complementary: we address correctness of low-level control, provide formal guarantees, and *check* rather than *assume* that runtime physics matches the model:

- The tools LTLMoP [9] and TuLiP [10] synthesize robot controls that satisfy a temporal logic specification. They excel at providing an intuitive user interface for specifying discrete planning problems, though discretization [11, 12] can be used to support continuous dynamics. We focus instead on providing the highest degree of confidence by proving safety in a theorem prover, including proofs of the dynamics and down to machine-code level.
- Controllers have been synthesized: *i*) from temporal logic specifications for linear systems [13], *ii*) for adaptive cruise control [14], tested in simulation and on hardware, and *iii*) from safety proofs [15] for switched systems using templates. These all assume model compliance and cannot ensure feedback controller correctness.

B. Offline Verification for Planning and Control

In contrast to online synthesis, offline verification can show safety in all uncountably many states. High-level models of the system under consideration can already be verified during the design phase of a project when changes are cheap. Much robotics verification work focuses on hybrid systems models; common approaches are reachability analysis [16] and theorem proving [6]. Both have been applied in case studies [7, 17] and experience shows that reachability typically provides more automation while theorem proving supports a powerful combination of rigorous foundations and establishes guarantees for unbounded time and space. Both approaches can be combined with monitoring:

- 1D straight-line motion was addressed both in dL [4] and with reachability analysis [18], but 1D uses simpler verification technology and is not suitable for real robots.
- Unbounded-time 2D obstacle avoidance and 1D liveness have also been proved in dL [7], and liveness has been

proved on paper [19]. Their controllers, like ours, are related to the Dynamic-Window [20] algorithm. Our novel results include 2D liveness, waypoint-following, and end-to-end correctness. While collision avoidance is simpler in prior work, their approach precludes liveness, which we proved. Prior dL efforts treat sensor errors explicitly, for which synthesis is subject of ongoing work [21]. In contrast, we integrate synthesized monitors with a simulation while keeping guarantees. To this end, we use a single tolerance for sensing/actuation error and deviation of real dynamics from the model. The limitation of this approach is that our guarantees do not explicitly incorporate sensor errors.

- A planner for ground vehicles was verified [8] in Isabelle. Their physics are close to ours, but feedback control and implementation correctness are not addressed.

C. Online Verification

Online/runtime verification provides a runtime safety net, but the correctness of the safety net itself is then critical to system safety. In contrast to offline verification, online verification cannot predict safety for infinitely many states.

- The basis of online verification is the Simplex [2] method, which uses a trusted monitor to decide between an untrusted controller and trusted fallback.
- The VeriPhy [4] toolchain for dL, which we use, combines offline and online verification to extend Simplex by ensuring the monitor is correct-by-construction, formally proving its safety, and maintaining those guarantees down to machine code implementations.
- Runtime monitoring has been combined with nonexhaustive model checking and evaluated in simulation [22]. Their relative strengths are in correctness of high-level event-handling logic and experimentally learning tolerances for the dynamics. Our relative strengths are use of a theorem-prover to show safety in *all* states, richer physical dynamics, and correct-by-construction monitors.
- Runtime reachability analysis has been used for Dubins-like car control [23], but runtime model compliance is not enforced and the reachability checker is trusted.

D. Simulation

Simulation is an essential part of evaluating models and designs. We used the AirSim [24] simulator for autonomous cars (originally for UAVs), because it comes with accurate physical and visual models out-of-the-box. Using these existing models provides a degree of independence in our evaluation.

In short, while verification of robotics receives frequent attention, few works have addressed rigorous end-to-end guarantees. We develop the first realistic system with formal end-to-end safety and liveness guarantees for 2D waypoint following, by generating a runtime monitor from a verified model. Crucially, we expect this runtime safety net can be applied to other Dubins-like system without redoing any proofs.

III. BACKGROUND: DIFFERENTIAL DYNAMIC LOGIC

We write our model as a *hybrid program* and use *differential dynamic logic* (dL) [5] to verify it. Hybrid programs express *hybrid systems* as programs containing differential equations (ODEs). They are particularly useful for verified robotics because they concisely describe both the control laws and kinematics of the system. Table I gives the syntax of hybrid programs and informally describes their semantics, wherein running a program α results in zero, one, or many different states. Detailed formal semantics are provided elsewhere [5]. Typical controllers use assignments $x := \theta$ to store the value of

TABLE I: Hybrid programs

Program	Means
$? \phi$	Results in current state if ϕ is true, no states if false.
$x := \theta$	Store value of expression θ in variable x .
$x := *$	Store arbitrary (real) number in variable x .
$x' = \theta \ \& \ \psi$	Evolve ODE $x' = \theta$ for any duration $t \geq 0$, with constraint formula ψ true throughout.
$\alpha; \beta$	Run α , then β in any resulting state(s).
$\alpha \cup \beta$	Choose between running α or β .
α^*	Repeats α n times, for any $n \in \mathbb{N}$.

(polynomial) expression θ in variable x , or assign an arbitrary value ($x := *$) and then test ($? \phi$) that the value satisfies some condition ϕ . Choice ($\alpha \cup \beta$) allows choosing between control laws, each of which may have (overlapping) tests ($? \phi$) saying when each law applies. Semicolons separate statements, so sequencing ($\alpha; \beta$) runs β after α , while loop α^* repeats α any arbitrary number of times. Many models follow the *control-plant loop* idiom (e.g., $\alpha_{1D} \equiv (1Dctrl; 1Dplant)^*$), where a discrete program $1Dctrl$ is followed by a continuous $1Dplant$ modeling physics, repeated in a loop ($*$). The $1Dplant$ is an ODE $x' = \theta \ \& \ \psi$, which evolves according to $x' = \theta$ for any duration such that ψ holds throughout. Before we develop a realistic 2D model in Sect. IV, we recall, in Example 1, a toy example, α_{1D} , of 1D motion with perfect speed control [4]. The controller can either **go** forward with some v such that $0 \leq v \leq V$ if we are far enough ($d \geq TV$) from the destination d , else it must **stop** by setting v to 0. The differential equation $d' = -v, t' = 1$ says the distance d continuously decreases proportional to velocity v , while time continuously elapses at rate 1. The constraint $t \leq T$ after $\&$ is a *time trigger*, saying that at most T seconds may elapse between control cycles. Note that we will show safety for any number of control cycles, and thus for *unbounded time*.

Example 1 (Simple 1D Idealized Driving).

$$1Dctrl \equiv go \cup stop \quad go \equiv ?d \geq TV; v := *; ?0 \leq v \leq V \\ stop \equiv v := 0 \quad 1Dplant \equiv t := 0; \{d' = -v, t' = 1 \ \& \ t \leq T\}$$

Formulas of dL are used to formalize program properties:

Definition 1 (dL formulas). *Formulas ϕ, ψ of dL consist of the following connectives:*

$$\phi, \psi ::= \phi \wedge \psi \mid \phi \vee \psi \mid \phi \rightarrow \psi \mid \neg \phi \mid \theta_1 \sim \theta_2 \\ \mid \forall x \phi \mid \exists x \phi \mid [\alpha] \phi \mid \langle \alpha \rangle \phi$$

where $\phi \wedge \psi$ holds when ϕ and ψ both hold, $\phi \vee \psi$ holds when either ϕ or ψ holds, $\phi \rightarrow \psi$ holds if ψ holds assuming

ϕ holds, $\neg\phi$ holds if ϕ does not, and $\forall x \phi$ and $\exists x \phi$ hold if ϕ holds for all or some value(s) of x , respectively. When we prove a theorem ϕ , all variables implicitly have a “for all” quantifier so, e.g., the safety and liveness theorems hold for *all* states and *all* values of system parameters. Formula $\theta_1 \sim \theta_2$ is shorthand for any comparison $\sim \in \{\leq, <, =, \neq, >, \geq\}$ where θ_i are real multivariate polynomials. The modalities $[\alpha]\phi$ and $\langle\alpha\rangle\phi$ say ϕ holds in *all* or *some* state(s) reached by executing hybrid program α respectively; they are used to express safety and liveness properties for our models.

Eq. 1 is a safety formula for α_{1D} : If the robot has not collided initially ($d \geq 0$) then the verified model will never collide no matter how many further control cycles are executed:

$$d \geq 0 \wedge V \geq 0 \wedge T \geq 0 \rightarrow [(1Dctrl; 1Dplant)^*] d \geq 0 \quad (1)$$

This toy example, which was previously used to demonstrate VeriPhy [4], misses out on many of the challenges essential to robotics: curved motion, acceleration, actuation disturbance, and goal-following all demand more sophisticated control conditions and invariants, which demand more sophisticated proof techniques. We take on these challenges in Sect. IV.

IV. GROUND ROBOT MODEL

This section introduces our 2D robot model in dL . This model is the heart of our verification effort: it lays out the definition of safety, assumptions on the controller, and assumptions on the plant. It will enable us in Sect. V to prove that these assumptions are strong enough to guarantee safety, then in Sect. VI to synthesize a monitor which functions as a runtime safety net, providing formal safety guarantees. The liveness proof of Sect. V complements safety and increases confidence in the model by showing our model is never so restrictive that it would force the robot to get stuck.

We use waypoint-following because it covers a wide variety of realistic scenarios, whereas collision avoidance is challenging to specify formally without making overly restrictive assumptions [7]. The implementation trusts waypoints from a planner. Feedback control is considered safe so long as it follows the waypoints within a fixed desired tolerance. The tolerance accounts for imperfect actuation and for discrepancies between the ideal dynamics and real dynamics of the implementation. The tolerance may also be increased to account for bounded sensor error, but the formal guarantees provided here are for perfect sensing.

Each waypoint is specified by coordinates (x, y) , a curvature k , and a speed limit $[v_l, v_h]$ for the robot’s velocity v . By convention, positive x points *forward*, and positive y points *left*. The curvature k yields circular arcs (when $k \neq 0$) and lines (when $k = 0$) as primitives. The addition of speed limits allows a plan to specify, for example, that the robot ought to slow down for a sharp curve or stop. The speed limits need only be met at the endpoint of the waypoints, which improves monitor compliance (in Sect. VI). Because realistic robots never follow a path perfectly, we bloat each arc to an annular section which is more easily followed. Our hybrid program α is again a time-triggered *control-plant* loop: $\alpha \equiv (ctrl; plant)^*$.

We use *relative* coordinates: the robot’s position is always the origin, from which perspective the waypoint “moves

toward” the vehicle. This simplifies proofs (fewer variables) and implementation (real sensors and actuators are vehicle-centric). Fig. 1 illustrates control scenarios for the system in relative coordinates. The robot is represented by $\blacktriangle$ where the triangle points in the robot’s forward direction. The control choice $k = 0$ drives waypoints that are straight ahead of the robot straight towards the robot. For waypoints initially to the left of the robot, a control choice $k > 0$ yields clockwise motion of the waypoint towards the robot. Conversely, for waypoints initially to the right, control choices with $k < 0$ yield counter-clockwise motion.

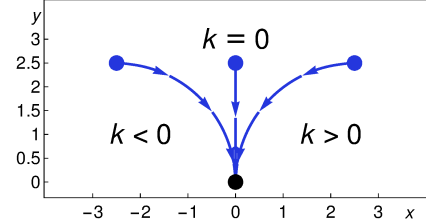


Fig. 1: Trajectories of dynamics for different choices of k .

The relative coordinate system and control choices for k are modeled by the ODE plant:

$$\begin{aligned} \text{plant} \equiv & t := 0; \{x' = v(ky - 1), y' = -vkx, v' = a, \\ & t' = 1 \ \& \ t \leq T \wedge v \geq 0\} \end{aligned}$$

Here, a is an input from the controller describing the acceleration with which the robot is to follow the arc of curvature k to waypoint (x, y) . In the equations for x', y' : *i*) The v factor models (x, y) moving at speed v , *ii*) The k, x, y factors model circular motion with curvature k . *iii*) The additional -1 term in the x' equation shifts the center of rotation to $(0, \frac{1}{k})$. The equations $v' = a$ and $t' = 1$ make acceleration the derivative of velocity and t stand for current time. The *domain constraint* $t \leq T \wedge v \geq 0$ after $\&$ says that the duration of one control cycle shall never exceed the timestep parameter $T > 0$ representing the maximum delay between control cycles and that the robot never drives in reverse.

Fig. 1 depicts curves that are exact solutions of *plant* where the robot exactly meets the waypoint. Because realistic robots cannot follow these curves exactly, the waypoint is bloated by a fixed ball of radius $\varepsilon > 0$, giving the robot some freedom in curve-following. We refer to this bloated waypoint as the *goal* for the robot. Fig. 2 illustrates a goal of size $\varepsilon = 1$ around the origin and several trajectories which pass through the goal.

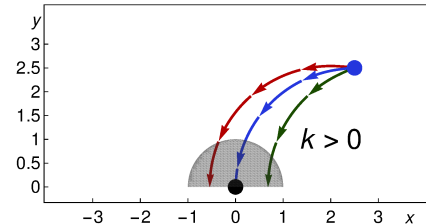


Fig. 2: Trajectories of *plant* for choices of $k < 0$ when $\varepsilon = 1$.

The controller’s task is to compute an acceleration a which slows down (or speeds up) soon enough that the speed limit

$v \in [v_l, v_h]$ is ensured by the time the robot reaches the goal. We allow the controller to exceed v_h temporarily as long as there is time to achieve $v \leq v_h$ before reaching the goal. This relaxation improves monitor compliance and lets the robot speed up more quickly, e.g., when it is far from the goal. The trade-off is that the proof becomes more challenging than it would be for a model that always enforces limits. The controller is written:

$$\begin{aligned} \text{ctrl} &\equiv (x, y) := *; [v_l, v_h] := *; k := *; ?\text{Feas}; \underbrace{a := *; ?\text{Go}}_{\text{ctrl}_a} \\ \text{Ann} &\equiv |k|\varepsilon \leq 1 \wedge \left| \frac{k(x^2 + y^2 - \varepsilon^2)}{2} - y \right| < \varepsilon \\ \text{Feas} &\equiv \text{Ann} \wedge x > 0 \wedge 0 \leq v_l < v_h \wedge \Delta T \leq v_h - v_l \wedge B T \leq v_h - v_l \\ \text{Go} &\equiv -B \leq a \leq A \wedge v + aT \geq 0 \\ &\wedge (v \leq v_h \wedge v + aT \leq v_h \vee \\ &\quad (1 + |k|\varepsilon)^2 \left(vT + \frac{a}{2}T^2 + \frac{(v+aT)^2 - v_h^2}{2B} \right) + \varepsilon \leq \|(x, y)\|_\infty) \\ &\wedge (v_l \leq v \wedge v_l \leq v + aT \vee \\ &\quad (1 + |k|\varepsilon)^2 \left(vT + \frac{a}{2}T^2 + \frac{v_l^2 - (v+aT)^2}{2A} \right) + \varepsilon \leq \|(x, y)\|_\infty) \end{aligned}$$

where the nondeterministic assignment $(x, y) := *$ chooses the next 2D waypoint, the assignment $[v_l, v_h] := *$ chooses the speed limit interval, and $k := *$ chooses any curvature. The subsequent *feasibility* test $?Feas$ checks whether or not the chosen waypoint, speed limit, and curvature are *physically* feasible in the current state under the *plant* dynamics (e.g., that there is enough remaining distance to get within the speed limit interval). We also simplify plans so that all waypoints satisfy $x > 0$ by subdividing any violating paths automatically. This simplifies the feedback controller and proofs.

In *Feas*, formula *Ann* says we are within the *annular section* (Fig. 3) ending at the waypoint (x, y) with the specified curvature k and width ε . A larger choice of ε yields more error tolerance in the sensed position and followed curvature at the cost of an enlarged goal region. Formula *Ann* also contains a simplifying assumption that the radius of the annulus is at least ε . *Feas* also says the speed limits are assumed distinct and large enough to not be crossed in one control cycle.

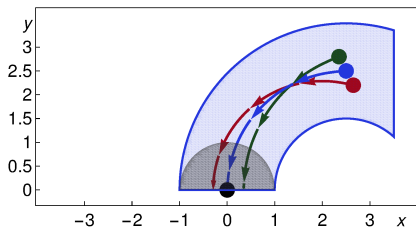


Fig. 3: Annular section through the (blue) waypoint (2.5, −3). Trajectories from the displaced green and red waypoints with slightly different curvatures remain within the annulus.

The *admissibility* test $?Go$ checks that the chosen a will take the robot to its goal with a safe speed limit, by *predicting future motion* of the robot. We illustrate this with the upper bound conditions. The bound will be satisfiable after one cycle

if either the chosen acceleration a already maintains speed limit bounds ($v \leq v_h \wedge v + aT \leq v_h$) or when there is enough distance left to restore the limits before reaching the goal. For straight line motion ($k = 0$), the required distance can simply be found by integrating acceleration and speed:

$$\underbrace{vT + \frac{a}{2}T^2}_{\text{distance at } T} + \underbrace{\frac{(v+aT)^2 - v_h^2}{2B}}_{\text{speed at } T} + \varepsilon \leq \|(x, y)\|_\infty$$

where $a \in [-B, A]$. The extra factor of $(1 + |k|\varepsilon)^2$ for curved motion accounts for the fact that an arc along the inner side of the annulus is shorter than one along the outside (Fig. 3).

The abbreviation ctrl_a names just the control code responsible for deciding *how* the waypoint is followed rather than *which* waypoint is followed. We show in Sect. VI that *Feas* and *Go* are easily satisfied by the implementation. We verify safety and liveness of the above model in Sect. V.

V. FORMAL SAFETY AND LIVENESS GUARANTEES

We now state the safety and liveness theorems in dL:

Theorem 1 (Safety). *The following dL formula is valid:*

$$\begin{aligned} &A > 0 \wedge B > 0 \wedge T > 0 \wedge \varepsilon > 0 \wedge J \rightarrow \\ &[(\text{ctrl}; \text{plant})^*](\|(x, y)\| \leq \varepsilon \rightarrow v \in [v_l, v_h]) \end{aligned}$$

where validity means the formula holds in all states and, e.g., all admissible waypoints.

The first four assumptions ($A > 0 \wedge \dots \wedge \varepsilon > 0$) are basic sign conditions on the symbolic constants used in the model. The final assumption, J , is the loop invariant. The full definition of J is deferred to Sect. V-A, but captures the fact that the robot never strays far from its path. We write $\|(x, y)\|$ for the Euclidean norm $\sqrt{x^2 + y^2}$ and consider the robot “close enough” to the waypoint when $\|(x, y)\| \leq \varepsilon$ for our chosen goal size ε . The theorem states that no matter which (admissible) control decisions are made, whenever the vehicle is in the goal region of size ε , it obeys the speed limit $v \in [v_l, v_h]$. While this provides a formal notion of safety, it does not prove that the robot can actually reach the goal, which is a *liveness* property:

Theorem 2 (Liveness). *The following dL formula is valid:*

$$\begin{aligned} &A > 0 \wedge B > 0 \wedge T > 0 \wedge \varepsilon > 0 \wedge J \rightarrow \\ &[(\text{ctrl}; \text{plant})^*](v > 0 \wedge x > 0 \rightarrow \\ &\quad \langle (\text{ctrl}_a; \text{plant})^* \rangle (\|(x, y)\| \leq \varepsilon \wedge v \in [v_l, v_h])) \end{aligned}$$

Under the same assumptions as Theorem 1, this theorem says that no matter how long the robot has been running ($[(\text{ctrl}; \text{plant})^*]$) already, if some simplifying assumptions still hold ($v > 0 \wedge x > 0$) the controller can be run ($\langle (\text{ctrl}_a; \text{plant})^* \rangle$) with admissible acceleration choices (ctrl_a) to reach the present goal ($\|(x, y)\| \leq \varepsilon$) within the desired speed limits ($v \in [v_l, v_h]$). The simplifying assumptions $v > 0 \wedge x > 0$ say the robot is still moving forward and the waypoint is still in the upper half-plane, i.e., it has not driven *past* the waypoint.

A. Proving Safety

We give the high-level insights here, such as invariants. We first prove the program stays in the region where J is true, then that J implies safety. To satisfy the safety condition, our program must maintain an invariant J that: *i*) the robot follows the plan closely and *ii*) it drives at speeds that let it achieve the speed limits in the remaining distance to the goal.

Predicate $\text{Lim}(v_1, v_2, a, \mathbf{x}, \mathbf{y}, \mathbf{k})$ says acceleration a can close velocity gap $v_1 - v_2$ before waypoint $(\mathbf{x}, \mathbf{y})$ reaches the origin:

$$\delta_{\text{Lim}}(v_1, v_2, a) \equiv (1 + |\mathbf{k}|\varepsilon)^2 \frac{v_1^2 - v_2^2}{2a}$$

$$\text{Lim}(v_1, v_2, a, \mathbf{x}, \mathbf{y}) \equiv v_1 \leq v_2 \vee \delta_{\text{Lim}}(v_1, v_2, a) + \varepsilon \leq \|(\mathbf{x}, \mathbf{y})\|_\infty$$

Here, $\delta_{\text{Lim}}(v_1, v_2, a)$ bounds distance needed to change velocity from v_1 to v_2 with a scaling factor $(1 + |\mathbf{k}|\varepsilon)^2$ for the tolerance incurred from the width of the annular section compared to an arc. This is not too conservative in practice: it is tightest (near 1) for small ε or small $|\mathbf{k}|$ and never exceeds 4. This suffices to define J , which says the speed limits can be achieved with maximum braking ($\mathbf{B}$) and acceleration ($\mathbf{A}$):

$$J \equiv \text{Ann} \wedge 0 \leq v_l < v_h \wedge \text{AT} \leq v_h - v_l \wedge \text{BT} \leq v_h - v_l \\ \wedge \text{Lim}(v_l, v, \mathbf{A}, \mathbf{x}, \mathbf{y}) \wedge \text{Lim}(v, v_h, \mathbf{B}, \mathbf{x}, \mathbf{y})$$

We proceed to prove Theorem 1 for waypoint-following safety. Per standard notation, the formula on the bottom (conclusion) is *valid* if all formulas on top (premises) are *valid*:

$$(\text{loop}) \quad \frac{P \rightarrow J \quad J \rightarrow Q \quad J \rightarrow [\alpha]J}{P \rightarrow [\alpha^*]Q}$$

The first two premisses, which prove automatically in KeYmaera X, say the invariant J holds initially and the postcondition follows from the invariant J , so staying inside J is safe.

The main task is the third premiss: loop body $\alpha \equiv \text{ctrl}; \text{plant}$ preserves the invariant J . The key conditions are **Feas** and **Go**; standard (automatic) dL proof steps reduce this to showing that J holds again after running the plant for time $t \leq T$:

$$(\text{auto}) \quad \frac{J, \text{Feas}, \text{Go} \rightarrow [\text{plant}]J}{J \rightarrow [\text{ctrl}; \text{plant}]J}$$

It remains to show the premiss of **auto**: **Feas** and **Go** imply that Ann , $\text{Lim}(v_h, v, \mathbf{B}, \mathbf{x}, \mathbf{y})$, and $\text{Lim}(v_l, v, \mathbf{A}, \mathbf{x}, \mathbf{y})$ *continue to hold* throughout the plant, which can be proved in dL using *differential induction* [5]. KeYmaera X helps prove invariants, but identifying them takes human ingenuity. For example, $\text{Lim}(v_h, v, \mathbf{B}, \mathbf{x}, \mathbf{y})$ is a natural choice, but because it is not inductive, we generalize it by hand in the full proof.

This completes the safety proof, showing that requirements **Feas** and **Go** guarantee that the robot obeys speed limits. We describe liveness next, which will reuse loop invariant J .

B. Proving Liveness

We show **Go** also allows the robot to reach the goal within speed limits $[v_l, v_h]$ by choosing correct acceleration a . The proof starts with the **loop** rule with invariant J as in Theorem 1, after which it remains to show:

$$J \wedge v > 0 \wedge x > 0 \rightarrow \langle (\text{ctrl}_a; \text{plant})^* \rangle (v \in [v_l, v_h] \wedge \|(\mathbf{x}, \mathbf{y})\| \leq \varepsilon)$$

Aside from the *invariants*, the key insight for liveness is a *progress function* which decreases as the waypoint approaches the origin [25]. There are multiple strategies to arrive at the goal within speed limit; for simplicity, we first enforce the speed limit $v \in [v_l, v_h]$ with appropriate acceleration choices and then maintain it until reaching the waypoint. This strategy splits the proof into the following two questions:

$$J \wedge v > 0 \wedge x > 0 \rightarrow \langle (\text{ctrl}_a; \text{plant})^* \rangle (v \in [v_l, v_h]) \quad (2)$$

$$v > 0 \wedge v \in [v_l, v_h] \wedge \text{Ann} \rightarrow \langle (\text{ctrl}_a; \text{plant})^* \rangle (v \in [v_l, v_h] \wedge \|(\mathbf{x}, \mathbf{y})\| \leq \varepsilon) \quad (3)$$

To prove Eq. 2, we pick the acceleration for ctrl_a in each of three situations: *i*) if the robot is too slow ($v < v_l$), it should speed up (pick $a = \mathbf{A}$), *ii*) if the speed is in the limits ($v \in [v_l, v_h]$) it maintains speed (pick $a = 0$), or *iii*) if it is too fast ($v > v_h$), it slows down (pick $a = -\mathbf{B}$). In all three cases, the invariant $v > 0 \wedge \text{Ann}$ is proved to be preserved throughout **plant**, so we soundly assume it in the proof of Eq. 3. The progress functions are case specific: e.g., in case *i*, the progress function $g \equiv v_l - v$, for the gap to the speed limit, decreases as the robot speeds up. Conversely, the progress function for case *iii* is $g \equiv v - v_h$.

Once the speed limit is achieved (case *ii*), the robot progresses toward the waypoint (Eq. 3) at a constant velocity ($a = 0$). The proof uses the progress function $g = x^2 + y^2 - \varepsilon^2$, i.e., the (squared) Euclidean distance to the goal region. The intuition for this progress function is shown in Fig. 4. The value of g is positive outside the goal region, strictly decreases along the trajectory, and is negative in the goal region.

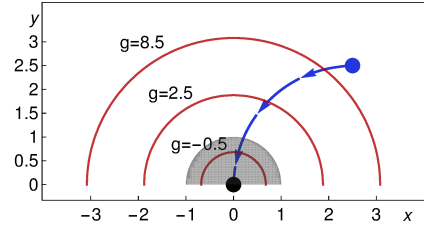


Fig. 4: Level sets of the progress function g (in red).

These are the crucial ingredients in the proof of liveness, which we have formally proved in KeYmaera X.

C. Proof Effort and Automation

User interactions are usually required for significant dL theorems like Theorem 1 (279 interactions) and Theorem 2 (589 interactions). User insight was mainly needed to choose invariants and progress functions for loops and ODEs. Most interactions are simplifications to help the automation. Automation handled most steps: 53,883 for safety, 225,607 for liveness. The final KeYmaera X proof scripts run automatically in 23s and 73s on a 2.4GHz i7 with 16GB memory.

VI. IMPLEMENTATION: SIMULATION

In this section, we fulfill the goal of extending verification to the level of simulation in AirSim. We use the VeriPhy

tool to synthesize an automatically-verified monitor containing both *controller* and *plant* monitor conditions. The controller monitor condition corresponds to **Feas** and **Go** in Sect. IV: any control decision satisfying these conditions is allowed and is guaranteed safe by Theorem 1, else the verified fallback is invoked. VeriPhy guarantees that the monitor is safe down to its machine code implementation, regardless what decisions are made by the external controller, *so long as* the plant monitor conditions are satisfied, which is the case so long as the differential invariant of Sect. V-A (the premiss of **auto**) holds for the sensed values. When the plant monitor conditions fail, safe braking is employed, albeit without the strong guarantee available in the other cases except with extra assumptions [3].

a) *High-Level Plans*: Our plan data structure is a graph (Fig. 5a) where waypoints are connected by lines and arcs, as in the dL model. The evaluation uses fixed plans (up to ≈ 80 segments), but our data structure also supports, e.g., nondeterministic plans (Fig. 5a, nodes B and F) and cyclic graphs for repeating missions, for the sake of flexibility.

b) *Feedback Control*: The high-level plan gives an ideal path to follow; the job of the low-level controller is to follow it within some tolerance. We give two classical feedback controllers: the *bang-bang* controller switches between hard-left and hard-right steering, while the *PD* scales to the discrepancy between current position and orientation vs. their target values. We compare the low-level controllers in AirSim, using a human operating AirSim as a baseline.

c) *AirSim Simulation*: We implemented our own plan representation and controllers in AirSim, ≈ 1100 lines of C++. We built the test environments (Fig. 5c, 5d, 5e) in Unreal Editor. Fig. 5b shows the AirSim car driving autonomously.

d) *Sensing and Actuation*: AirSim does not explicitly simulate sensing and actuation error, but some implementation details of the kinematics are unknown, so actuation error must be accounted for in practice. The tolerance ε does not include sensing error, but does account for deviation of the AirSim kinematics from ideal Dubins. Thanks to the proofs and sensors, errors do not accumulate: if actuation is imperfect, the deviation is detected by the sensors and feedback control counters the deviation as usual. If our monitor conditions are applied in systems where sensors accumulate drift over time, it would not obviate the need to account for those drifts.

e) *Results*: We assess monitor compliance and safety of each controller. We assess liveness indirectly by checking how *quickly* the goal is reached. We assess compliance and safety directly: a successful simulation should comply with the monitor conditions (especially the safety-critical plant monitor conditions) a large majority of the time and have no safety violations. Our three AirSim environments are shown in Fig. 5. These environments cover medium turns at medium speed (Fig. 5c), tight turns at low speed (Fig. 5d), and wide curves at high speed (Fig. 5e). We simulated bang-bang and PD controllers of different speeds driving each environment (Table II), with amateur human pilots as a baseline.

The car completes every environment, except “Clover” where bang-bang control fails to track long curves. As promised, there were no safety violations. The controller monitor condition has few failures. The plant monitor condition

fails more often, but rarely enough that the car completes the drive. The plant monitor failed more since our bloated ideal dynamics are simpler than the AirSim physics. In general, the failure rate increases the more physics differs from Dubins.

We ran the tests with $\varepsilon = 1\text{m}$, which was small enough to stress-test the controllers. For our purposes, the exact value of ε is less important than the fact that safety is guaranteed for all values of ε . The bang-bang controller exhibits tracking error at speed and so did not complete the Clover track. The slower PD controller (PD1, in bold) had the best (lowest) overall error rate. The human and the remaining controllers had high plant failure rate on the “Clover” level due to tracking error. PD control (particularly PD3) had speed and monitor failure rates competitive with the human baseline. The bang-bang controller’s rough steering increased its plant failure rate.

While complete model compliance is a challenge, well-tuned controllers came close in all environments, *even though* the model is simple. The crucial insight is that the bloated model allows realistic imperfections in actuation, and that the *untrusted controller* makes steering and acceleration choices that satisfy its monitor condition. Most of the time, formal guarantees apply because both monitor conditions are satisfied. The plant’s monitor condition detects the few cases where guarantees do not apply, engages the fallback action, and then returns to normal control without any actual safety violation. The monitor furthermore helps us or any other developer identify and reduce the remaining non-compliant cases.

VII. CONCLUSIONS AND LESSONS

We cast a formally verified safety net that provides end-to-end verification guarantees for 2D Dubins waypoint-following. We developed a dL model, proved it safe and live in KeYmaera X, then synthesized a verified monitor with ModelPlex, and synthesized verified machine code with VeriPhy. The resulting safety net ensures safety even with unverified robot software so long as *plant* assumptions hold and collision-free plans are provided. We simulate the system in AirSim with several controllers; our aim was not to innovate in controller design, but to show that monitors generated from dL proofs can be applied in realistic scenarios, thanks largely to the use of verified tolerances in the model and proof. The evaluation showed that our simple tolerance-based model did not hinder applicability, because even realistic simulations look like Dubins at a distance. Our simple model greatly eased formal verification. Just as we improved on prior models [4], future work can verify more sophisticated models to improve compliance or reduce the tolerance ε .

The second major direction for future work is to apply VeriPhy on real robot hardware, and as a development aid for novel robots. Because the KeYmaera X proofs are significantly more complex than the sketches presented here (see Sect. V-C), VeriPhy’s reusability is essential to make it practical as a development tool. We are presently in the process of reusing our current monitors as-is on a hardware platform that follows Dubins paths without changing the proofs, and subsequently intend to pursue more challenging motion scenarios that violate our model’s assumption.

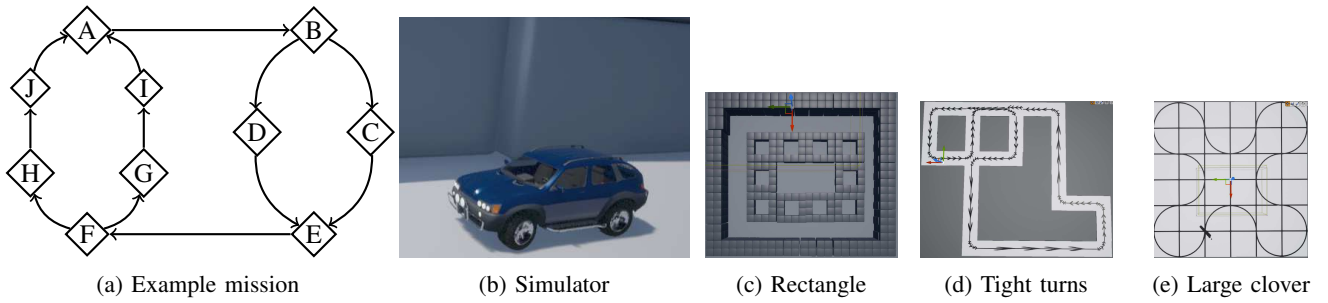


Fig. 5: Implementation and environments built in AirSim

TABLE II: Average speed, Monitor failure rates, for AirSim and human driver in each environment

World	Avg. Speed (m/s)					Ctrl Fail.					Plant Fail.				
	BB	PD1	PD2	PD3	Human	BB	PD1	PD2	PD3	Human	BB	PD1	PD2	PD3	Human
Rect	4.3	6.32	7.16	12.6	9.92	0.5%	0.1%	0.1%	0.19%	1.14%	36.8%	8.23%	8.5%	14%	41.3%
Turns	3.78	3.95	4.43	4.69	9.66	1.0%	1.0%	1.1%	4.7%	3.61%	18.6%	3.95%	6.8%	11%	21.1%
Clover	X	29.5	29.5	29.5	28.9	X	0.2%	0.2%	0.19%	0.29%	X	66%	66%	66%	48%

REFERENCES

- [1] T. A. Henzinger, "The theory of hybrid automata," in *LICS*. IEEE, 1996.
- [2] D. Seto, B. Krogh, L. Sha, and A. Chutinan, "The Simplex architecture for safe online control system upgrades," in *ACC*, 1998.
- [3] S. Mitsch and A. Platzer, "ModelPlex: verified runtime validation of verified cyber-physical system models," *Form. Methods Syst. Des.*, 2016.
- [4] B. Bohrer, Y. K. Tan, S. Mitsch, M. O. Myreen, and A. Platzer, "VeriPhy: verified controller executables from verified cyber-physical system models," in *PLDI*, 2018.
- [5] A. Platzer, *Logical Foundations of Cyber-Physical Systems*. Springer, 2018.
- [6] N. Fulton, S. Mitsch, J. Quesel, M. Völz, and A. Platzer, "KeYmaera X: an axiomatic tactical theorem prover for hybrid systems," in *CADE*, 2015.
- [7] S. Mitsch, K. Ghorbal, D. Vogelbacher, and A. Platzer, "Formal verification of obstacle avoidance and navigation of ground robots," *I. J. Robotics Res.*, 2017.
- [8] A. Rizaldi, F. Immler, B. Schürmann, and M. Althoff, "A formally verified motion planner for autonomous vehicles," in *ATVA*, ser. LNCS, vol. 11138, 2018.
- [9] C. Finucane, G. Jing, and H. Kress-Gazit, "LTLMoP: Experimenting with language, temporal logic and robot control," in *IROS*. IEEE, 2010.
- [10] I. Filippidis, S. Dathathri, S. C. Livingston, N. Ozay, and R. M. Murray, "Control design for hybrid systems with TuLiP: The temporal logic planning toolbox," in *CCA*. IEEE, 2016.
- [11] J. Liu, N. Ozay, U. Topcu, and R. M. Murray, "Synthesis of reactive switching protocols from temporal logic specifications," *IEEE Trans. Automat. Contr.*, 2013.
- [12] G. E. Fainekos, A. Girard, H. Kress-Gazit, and G. J. Pappas, "Temporal logic motion planning for dynamic robots," *Automatica*, 2009.
- [13] M. Kloetzer and C. Belta, "A fully automated framework for control of linear systems from temporal logic specifications," *IEEE Trans. Automat. Contr.*, 2008.
- [14] P. Nilsson, O. Hussien, A. Balkan, Y. Chen, A. D. Ames, J. W. Grizzle, N. Ozay, H. Peng, and P. Tabuada, "Correct-by-construction adaptive cruise control: Two approaches," *IEEE Trans. Contr. Sys. Techn.*, 2016.
- [15] A. Taly and A. Tiwari, "Switching logic synthesis for reachability," in *EMSOFT*. ACM, 2010.
- [16] G. Frehse, C. L. Guernic, A. Donzé, S. Cotton, R. Ray, O. Lebeltel, R. Ripado, A. Girard, T. Dang, and O. Maler, "SpaceX: Scalable verification of hybrid systems," in *CAV*, ser. LNCS, vol. 6806. Springer, 2011.
- [17] C. Tomlin, I. Mitchell, and R. Ghosh, "Safety verification of conflict resolution manoeuvres," *IEEE Trans. Int. Trans. Sys.*, 2001.
- [18] X. Chen, S. Schupp, I. B. Makhoul, E. Ábrahám, G. Frehse, and S. Kowalewski, "A benchmark suite for hybrid systems reachability analysis," in *NFM*, ser. LNCS, vol. 9058. Springer, 2015.
- [19] B. Martin, K. Ghorbal, E. Goubault, and S. Putot, "Formal verification of station keeping maneuvers for a planar autonomous hybrid system," in *FVAV@iFM*, 2017.
- [20] D. Fox, W. Burgard, and S. Thrun, "The dynamic window approach to collision avoidance," *IEEE Robot. Automat. Mag.*, 1997.
- [21] S. Mitsch and A. Platzer, "Verified runtime validation for partially observable hybrid systems," *CoRR*, vol. abs/1811.06502, 2018.
- [22] A. Desai, T. Dreossi, and S. A. Seshia, "Combining model checking and runtime verification for safe robotics," in *RV*, ser. LNCS, vol. 10548. Springer, 2017.
- [23] M. Althoff and J. M. Dolan, "Online verification of automated road vehicles using reachability analysis," *IEEE Trans. Robot.*, 2014.
- [24] S. Shah, D. Dey, C. Lovett, and A. Kapoor, "AirSim: High-fidelity visual and physical simulation for autonomous vehicles," in *FSR*. Springer, 2017.
- [25] A. Sogokon and P. B. Jackson, "Direct formal verification of liveness properties in continuous and hybrid dynamical systems," in *FM*, ser. LNCS, vol. 9109. Springer, 2015.