

Agile But Safe: Learning Collision-Free High-Speed Legged Locomotion

Tairan He^{1†} Chong Zhang^{2†} Wenli Xiao¹ Guanqi He¹ Changliu Liu¹ Guanya Shi¹

¹Carnegie Mellon University ²ETH Zürich [†]Equal Contributions

Page: <https://agile-but-safe.github.io> Code: <https://github.com/LeCAR-Lab/ABS>



Fig. 1. Our proposed framework **ABS** demonstrates agile *and* collision-free locomotion capabilities, where the robot, with fully onboard computation and sensing, can safely navigate through cluttered environments and rapidly react to diverse and dynamic obstacles, both indoors and outdoors. **ABS** involves a dual-policy setup: **green** lines at the bottom indicate the **agile policy** taking control, and **red** lines indicate the **recovery policy** in operation. The **agile policy** enables the robot to run fast amidst obstacles, and the **recovery policy** saves the robot from risky cases where the **agile policy** might fail. **Subfigures:** (a) The robot dodges a swinging human leg. (b) The **agile policy** enables the robot to run at a peak speed of 3.1 m/s. (c) The robot dodges a moving stroller during high-speed locomotion. (d) The robot dodges a moving human in snowy terrain. (e) The robot safely navigates in a hall with both static and dynamic obstacles, with an average speed of 2.1 m/s and a peak speed of 2.9 m/s. (f) The robot avoids obstacles and moving humans in a dim corridor, with an average speed of 1.5 m/s and a peak speed of 2.5 m/s. (g) The robot, running outdoors at an average speed of 2.3 m/s and a peak speed of 3.0 m/s, avoids both moving and static trash bins and climbs up a grassy slope. **Videos:** see the website.

Abstract—Legged robots navigating cluttered environments must be jointly *agile* for efficient task execution and *safe* to avoid collisions with obstacles or humans. Existing studies either develop conservative controllers (< 1.0 m/s) to ensure safety, or focus on agility without considering potentially fatal collisions. This paper introduces Agile But Safe (ABS), a learning-based control framework that enables agile and collision-free locomotion for quadrupedal robots. ABS involves an agile policy to execute agile motor skills amidst obstacles and a recovery policy to prevent failures, collaboratively achieving high-speed and collision-free navigation. The policy switch in ABS is governed by a learned control-theoretic reach-avoid value network, which also guides the recovery policy as an objective function, thereby safeguarding the robot in a closed loop. The training process involves the learning of the agile policy, the reach-avoid value network, the recovery policy, and an exteroception representation network, all in simulation. These trained modules can be directly deployed in the real world with onboard sensing and computation, leading to high-speed and collision-free navigation in confined indoor and outdoor spaces with both static and dynamic obstacles (Figure 1).

I. INTRODUCTION

Agile locomotion of legged robots in cluttered environments presents a non-trivial challenge due to the inherent trade-

off between agility and safety, and is crucial for real-world applications that require both robustness and efficiency, such as search and rescue [68], disaster response [63], and police robotics [71]. Existing works typically exhibit limited agility (velocity < 1 m/s) to ensure safety [38, 17, 7, 23, 28, 77, 12, 49, 44, 83], or focus solely on maximizing agility without considering safety in navigation scenarios [48, 58]. Our work distinguishes itself by achieving high-speed (max velocity > 3 m/s), collision-free quadrupedal locomotion in cluttered environments.

The agility limitations in existing works in the navigation domain stem from varied factors. Regarding the formulation, some decouple locomotion and navigational planning into two subtasks and build hierarchical systems [38, 17, 7, 28, 49, 83]. Such decoupling not only constrains the controller from the optimal solution [61], but also results in conservative behaviors to ensure safety, thereby limiting the system from fully unleashing the locomotion agility. This work, instead, learns end-to-end controllers that directly output joint-level actions for collision-free locomotion to reach specified goal positions. Our approach is inspired by recent works [76, 54, 82, 29] where

robots learn end-to-end controllers to overcome challenging terrains by integrating locomotion with navigation.

Regarding the controller, some works employ model-based methods with simplified models, such as model predictive control (MPC) and barrier functions, for guaranteed safety [23, 12, 44]. The model mismatch and potential constraint violations such as slippage, together with the online computational burden, limit these controllers from agile motions and stable deployment in the wild [23, 34, 41]. On the other hand, recent progress of model-free reinforcement learning (RL) in legged locomotion has demonstrated remarkable agile motor skills that model-based controllers have not achieved [48, 58, 29, 41, 50, 36, 86, 46, 79, 11], although potentially unsafe in cluttered environments. We harness the flexibility and agility of model-free RL and further safeguard it using control-theoretic tools.

Named ABS, our framework goes beyond a single RL policy. First, we have a model-free perceptive agile policy that incorporates collision avoidance into locomotion, as presented in Section IV, enabling our Go1 robot to achieve peak speeds up to 3.1 m/s while being aware of collisions. However, the RL policy does not guarantee safety, so we safeguard the robot with another recovery policy (see Section VI) when the agile policy may fail. To decide which policy to take control, we use a policy-conditioned reach-avoid (RA) value network to quantify the risk level of the agile policy. This is inspired by [30] where model-free RA values can be efficiently learned based on the Hamilton-Jacobi reachability theory [3]. The RA value network is trained by a discounted RA Bellman equation, with data collected by the learned agile policy in simulation. Beyond being a threshold, the differentiable RA value network also provides gradient information to guide the recovery policy, thus closing the loop, which will be further presented in Section V.

To get collision avoidance behaviors that can generalize in different scenarios, we use a low-dimensional exteroceptive feature for policy and RA value training: the traveling distances of several rays cast from the robot to obstacles. In order to achieve this, we additionally train an exteroception representation (or ray-prediction) network with simulated data, which maps depth images to ray distances as detailed in Section VII. By doing so, we achieve robust collision avoidance in high-speed locomotion with onboard sensing and computation.

Briefly, we identify our contributions as follows:

- 1) A perceptive agile policy for obstacle avoidance in high-speed locomotion with novel training methods.
- 2) A novel control-theoretic data-driven method for RA value estimation conditioned on the learned agile policy.
- 3) A dual-policy setup where an agile policy and a recovery policy collaborates for high-speed collision-free locomotion, and the RA values govern the policy switch and guide the recovery policy.
- 4) An exteroception representation network that predicts low-dimensional obstacle information for generalizable collision avoidance capability.
- 5) Validation of ABS's superior safety and state-of-the-art agility amidst obstacles indoors and outdoors (Figure 1).

II. RELATED WORKS

A. Agile Legged Locomotion

Model-based methods such as MPC use simplified models and handcrafted gaits to enable dynamic legged locomotion [5, 14, 15, 25, 35, 26]. Despite their impressive performance in simulation and under laboratory conditions, these methods struggle in the wild due to model mismatch and unexpected slippage [34, 41]. The online computational burden also limits perceptive model-based controllers from agile motions [12].

Recently, RL-based controllers have shown promising results for robust locomotion [76, 41, 50, 24] and agile motor skills including high-speed running [32, 48, 58], challenging terrain traversal [82, 29, 36, 86, 11], jumping [42, 79], and fall recovery [76, 46, 72, 81]. However, existing works on agile locomotion mostly study how to achieve fast speeds for racing or skillful motions to overcome challenging terrains. In cluttered environments, these methods necessitate a high-level navigation module for collision avoidance, which is typically conservative and greatly constrains the motion far below the motor limit [73, 83]. In contrast, this paper studies agile collision avoidance for versatile navigation.

B. Legged Collision Avoidance

Classical methods tackle collision avoidance in legged robots with collision-free motion planning [38, 17, 7] in the configuration space without considering the robot dynamics, leading to slow and statically stable gaits. MPC-based methods [23, 12, 64, 44] integrate planning and control by treating distances to obstacles as optimization constraints. However, they suffer from the aforementioned drawbacks of model-based controllers and slow movements (velocity < 0.5 m/s).

Learning-based methods are another choice. Some existing works [28, 57, 37, 69, 83, 80] train RL-based policies that output twist commands to be tracked by the locomotion controller, while the velocity commands are bounded by 1 m/s to ensure safety. However, the decoupling of navigation planning and locomotion control makes high-speed locomotion risky, as the high-level planner is unaware of the low-level tracking error.

Yang et al. [77] instead provides an end-to-end RL-based solution that maps depth images and proprioceptive data directly to joint actions, but the robot can only walk forward and the velocity is limited to ~ 0.4 m/s. In contrast, our work deploys an end-to-end agile policy for omnidirectional rapid locomotion with collision avoidance, and safeguards the robot with RA values and a recovery policy. To the best of our knowledge, our work is the first to validate collision-free quadruped locomotion with maximum velocity up to 3.1 m/s. Even in tight space with dynamic adversarial obstacles, our system can still reach a peak velocity of 2.5 m/s and an average speed of 1.5 m/s (Figure 1 (f)).

C. Safe Reinforcement Learning

There are two main categories of methods to perform safe RL [85]: 1) *end-to-end* methods and 2) *hierarchical* methods. Lagrangian-based methods [53, 4, 43, 65] are the most representative *end-to-end* safe RL methods that solve a primal-dual

optimization problem to satisfy safety constraint where the Lagrange multipliers can be optimized along with the policy parameters. However, the constraint is often enforced before convergence, hindering exploration and lowering returns [53].

Hierarchical safe RL methods safeguard unsafe RL actions using structures of underlying dynamics [13, 78, 75] and control-theoretic safety certificates [10, 84, 51]. These methods typically build on the assumptions of available dynamics or safety certificate functions before learning, which heavily limits the scalability to high-dimensional complex systems. Some recent works learn safety prediction networks (or safety critics) and safety backup policies to safeguard RL when the safety critics indicate the nominal policy is unsafe [66, 31]. Nevertheless, these frameworks lack interplay between safety critics and backup policies, relying on the demanding assumption that the backup policy can restore safety without explicit optimization to satisfy the safety critics.

Our approach aligns with the *hierarchical* methods, yet it stands out with a distinctive strategy. We focus on estimating the reach-avoid values of the agile policy and feed the reach-avoid values' gradient information back into the system to guide the recovery policy within a closed loop. This innovative approach enables a dynamic and adaptive recovery process. Notably, all our modules are trained in simulation using a model-free approach, enhancing the generalizability and scalability of our method.

D. Reach-Avoid Problems and Hamilton-Jacobi Analysis

Reach-avoid (RA) problems involve navigating a system to reach a target while avoiding certain undesirable states. Hamilton-Jacobi (HJ) reachability analysis [3] solves this problem by analyzing the associated Hamilton-Jacobi partial differential equation, which provides a set of states that the system must stay out of in order to remain safe.

HJ reachability analysis faces computational challenges, which escalate exponentially with the system's dimensionality [9]. Recent learning-based methods [2] try to scale HJ reachability analysis to high-dimensional systems by learning value networks that satisfy the associated HJ partial differential equations and constraints. However, they still require explicit system Hamiltonian expression before learning.

Our method builds on another line of works [21, 30] that leverage contraction properties to derive a time-discounted reach-avoid Bellman equation. However, contrary to previous works that learn policy-agnostic RA values during RL training, we instead learn a policy-conditioned RA value network. This not only reduces the computational burden by avoiding the identifiability issue of the global RA set but also best suits our trained agile policy. Similarly, a concurrent work on neural control barrier functions [60] also applies policy-conditioned safety filters to shield complex systems with reduced complexity.

TABLE I
IMPORTANT SYMBOLS AND ABBREVIATIONS

Symbol	Meaning
t	Time step, converted to time in calculation
$s \in \mathcal{S}$	State
$a \in \mathcal{A}$	Action
$o \in \mathcal{O}$	Observation
T	Time horizon or episode length
γ_{RL}	Discount factor for reinforcement learning
γ_{RA}	Discount factor in reach-avoid Bellman equation
$V_{\text{threshold}}$	Reach-avoid value threshold equal to $-\epsilon, \epsilon > 0$
$\zeta(\cdot)$	Function indicating failures
$l(\cdot)$	Function indicating reaching the target
$V_{\text{RA}^*}^{\pi}(\cdot)$	Ground-truth reach-avoid values conditioned on policy π
$V_{\text{RA}}^{\pi}(\cdot)$	Ground-truth discounted reach-avoid values conditioned on policy π
$\hat{V}(\cdot)$	Neural network for discounted policy-conditioned reach-avoid value approximation
π^{Agile}	Agile policy
π^{Recovery}	Recovery policy
v	Linear velocity in the base frame
ω	Angular velocity in the base frame
c_f	Foot contact statuses
g	Normalized projected gravity in the base frame
$q, \dot{q}, \ddot{q}$	Joint positions, velocities, and accelerations
τ	Joint torques
R	Logarithm of ray distances
d_{goal}	Distance from the robot to the goal
G^c	Goal command
v^c	Linear velocity command in the base frame
ω^c	Angular velocity command in the base frame
tw^c	Twist command
$\text{ReLU}(\cdot)$	Function clipping negative values to zero [22]
Abbreviation	Full Form
ABS	Agile but safe
RA	Reach-avoid
RL	Reinforcement learning
MPC	Model predictive control
MLP	Multilayer perceptron

III. OVERVIEW AND PRELIMINARIES

A. Nomenclature

We present important symbols and abbreviations that are used across this paper in Table I for reference.

B. Problem Formulation

1) *Dynamics*: Let $s_t \in \mathcal{S} \subset \mathbb{R}^{n_s}$ be the state at time step t , where n_s is the dimension of the state space $\mathcal{S}$; $a_t \in \mathcal{A} \subset \mathbb{R}^{n_a}$ be the control input at time step t , where n_a is the dimension of the action space $\mathcal{A}$. The system dynamics are defined as:

$$s_{t+1} = f(s_t, a_t), \quad (1)$$

where $f : \mathcal{S} \times \mathcal{A} \rightarrow \mathcal{S}$ is a function that maps the current robot state and control to the next state. For simplicity, this paper considers deterministic dynamics that can be without an analytical form. We denote the robot observations from proprioception and/or exteroception as $o_t = h(s_t)$ where $h : \mathcal{S} \rightarrow \mathcal{O}$ is the sensor mapping. Detailed observation and action space of the agile policy and recovery policy will be introduced in Section IV and Section VI.

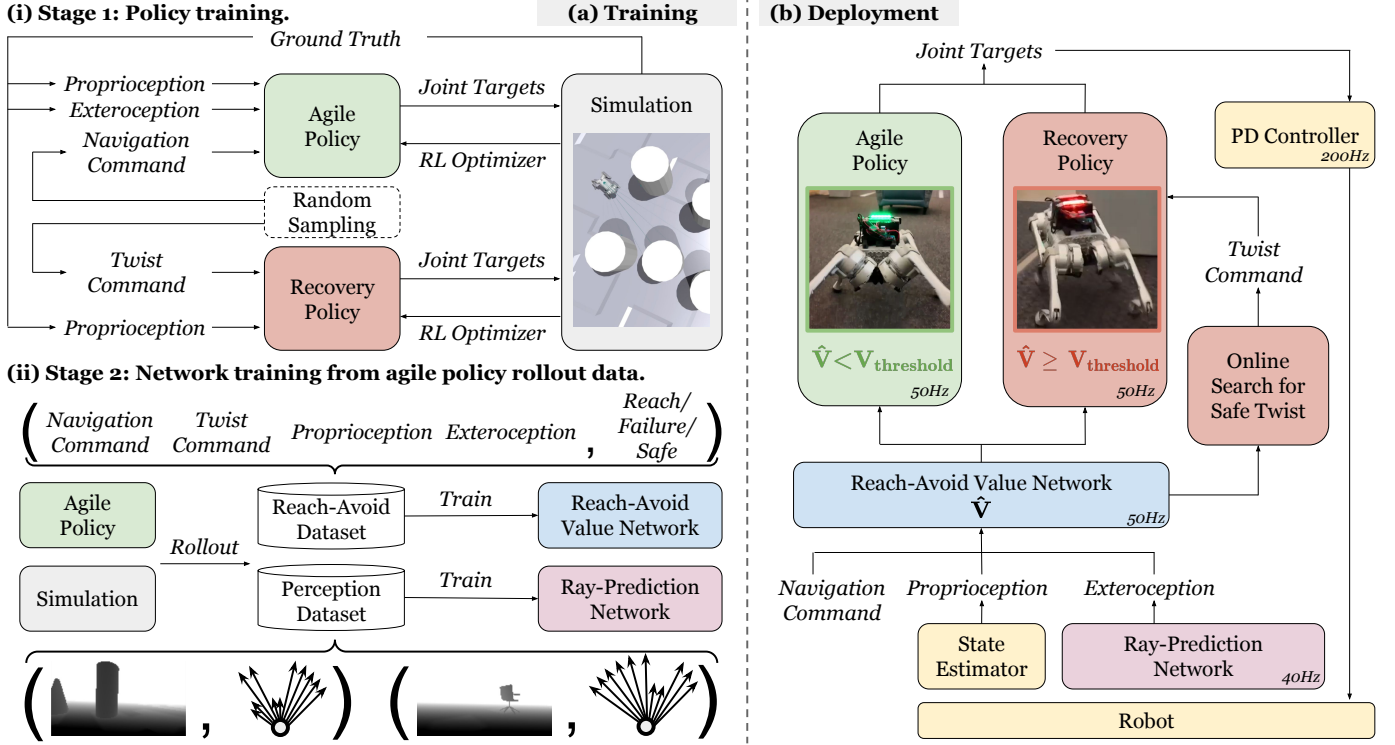


Fig. 2. Overview of ABS: (a) There are four trained modules within the ABS framework: 1) **Agile Policy** (introduced in Section IV) is trained to achieve the maximum agility amidst obstacles; 2) **Reach-Avoid Value Network** (introduced in Section V) is trained to predict the RA values conditioned on the agile policy as safety indicators; 3) **Recovery Policy** (introduced in Section VI) is trained to track desired twist commands (2D linear velocity v_x^c, v_y^c and yaw angular velocity ω_z^c) that lower the RA values; 4) **Ray-Prediction Network** (introduced in Section VII) is trained to predict ray distances as the policies' exteroceptive inputs given depth images. (b) Illustration of the ABS deployment architecture. The dual policy setup switches between the *agile policy* and the *recovery policy* based on the estimated $\hat{V}$ from the RA value network: 1) if $\hat{V} < V_{\text{threshold}}$, the *agile policy* is activated to navigate amidst obstacles; 2) if $\hat{V} \geq V_{\text{threshold}}$, the *recovery policy* is activated to track twist commands that lower the RA values via constrained optimization.

2) *Goal and Policy*: Goal-conditioned reinforcement learning [45] learns to reach goal states $G \in \Gamma$ via a goal-conditioned policy $\pi : \mathcal{O} \times \Gamma \rightarrow \mathcal{A}$. With the reward function $r : \mathcal{S} \times \mathcal{A} \times \Gamma \rightarrow \mathbb{R}$ and the discount factor γ_{RL} . The policy is learned to maximize the expected cumulative return over the goal distribution p_G :

$$J(\pi) = \mathbb{E}_{a_t \sim \pi(\cdot | o_t, G), G \sim p_G} \left[\sum_t \gamma_{\text{RL}}^t r(s_t, a_t, G) \right]. \quad (2)$$

3) *Failure Set, Target Set and Reach-Avoid Set*: We denote the *failure set* $\mathcal{F} \subseteq \mathcal{S}$ as unsafe states (e.g., collision) where the robot is not allowed to enter. The *failure set* can be represented by the zero-sublevel set of a Lipschitz-continuous function $\zeta : \mathcal{S} \rightarrow \mathbb{R}$, i.e., $s \in \mathcal{F} \Leftrightarrow \zeta(s) > 0$. The *target set* $\Theta \subset \mathcal{S}$ is defined as desired states (i.e., goal states). Similarly, the *target set* can be represented by the zero-sublevel set of a Lipschitz-continuous function $l : \mathcal{S} \rightarrow \mathbb{R}$, i.e., $s \in \Theta \Leftrightarrow l(s) \leq 0$. We denote $\xi_{s_t}^\pi(\cdot)$ as the future trajectory rollout from state s_t ($\xi_{s_t}^\pi(0) = s_t$) using policy π up to s_T . The *reach-avoid set* conditioned on policy π is defined as

$$\mathcal{RA}^\pi(\Theta; \mathcal{F}) := \{s_t \in \mathcal{S} \mid \xi_{s_t}^\pi(T-t) \in \Theta \wedge \forall t' \in [0, T-t], \xi_{s_t}^\pi(t') \notin \mathcal{F}\}, \quad (3)$$

which represents the set of states governed by policy π capable of leading the system to Θ while consistently avoiding $\mathcal{F}$ in

all prior timesteps.

4) *Reach-Avoid Value and Time-Discounted Reach-Avoid Bellman Equation*: We define policy-conditioned reach-avoid values as: $V_{\text{RA}}^\pi(s) \leq 0 \Leftrightarrow s \in \mathcal{RA}^\pi(\Theta; \mathcal{F})$. Following the proof (Appendix A in [30]), it can be easily extended that value function $V_{\text{RA}}^\pi(s)$ satisfies the fixed-point reach-avoid Bellman equation (our policy-conditioned value function is a special case of the general value function):

$$V_{\text{RA}}^\pi(s) = \max \left\{ \zeta(s), \min \{ l(s), V_{\text{RA}}^\pi(f(s, \pi(s))) \} \right\}. \quad (4)$$

However, there is no assurance that Equation (4) will result in a contraction in the space of value functions. To make it accessible to data-driven approximation, we leverage time-discounted reach-avoid Bellman equation [30] to make a contraction on the discounted policy-conditioned reach-avoid values $V_{\text{RA}}^\pi(s)$ defined as

$$V_{\text{RA}}^\pi(s) = \gamma_{\text{RA}} \max \left\{ \zeta(s), \min \{ l(s), V_{\text{RA}}^\pi(f(s, \pi(s))) \} \right\} + (1 - \gamma_{\text{RA}}) \max \{ l(s), \zeta(s) \}. \quad (5)$$

Following [30], it can be shown that $V_{\text{RA}}^\pi(s)$ is always an under-approximation of $V_{\text{RA}}^\pi(s)$ for $\gamma_{\text{RA}} \in [0, 1)$, and $V_{\text{RA}}^\pi(s)$ converges to $V_{\text{RA}}^\pi(s)$ as γ_{RA} approaches 1. Note that

the under-approximation of $V_{RA}^\pi(s)$ to $V_{RA}^{\pi^*}(s)$ means that $V_{RA}^\pi(s) \leq 0 \Rightarrow s \in \mathcal{RA}^\pi(\Theta; \mathcal{F})$, which enables that shielding methods on thresholds of $V_{RA}^\pi(s)$ could make the system stay in the control-theoretic *reach-avoid set* $\mathcal{RA}^\pi(\Theta; \mathcal{F})$.

C. System Structure

As shown in Figure 2, our proposed ABS framework involves a dual-policy setup where the agile policy π^{Agile} and the recovery policy π^{Recovery} work together to enable agile and safe locomotion skills. The agile policy performs agile motor skills (up to 3.1 m/s on Unitree Go1) to navigate the robot based on goal commands (target 2D positions and headings) with basic collision-avoidance ability (see also Section IV). The recovery policy is responsible for safeguarding the agile policy by rapidly tracking twist commands (2D linear velocity v_x^c, v_y^c and yaw rate ω_z^c) that can avoid collisions (see also Section V-C and Section VI). Both policies output joint targets that are tracked by a PD controller.

During deployment, the policy switch is governed by RA values conditioned on the agile policy, estimated using a neural network $\hat{V}$ (see also Section V). With a safety threshold $V_{\text{threshold}} = -\epsilon$ where ϵ is a small positive number, we have:

- If $\hat{V} \geq V_{\text{threshold}}$, we search for a twist command that drives the robot closer to the goal while maintaining safety based on $\hat{V}$ (see also Equation (21)). The recovery policy takes control and tracks the searched twist command.
- If $\hat{V} < V_{\text{threshold}}$, the agile policy takes control.

We expect the system to activate the agile policy in most time, and use the recovery policy as a safeguard in risky situations until it is safe again for the agile policy, *i.e.*, $\hat{V} < V_{\text{threshold}}$.

For collision avoidance, both the agile policy and the RA value networks need exteroceptive inputs. Inspired by [29, 16, 1], we choose to use a low-dimensional exteroception representation: the distances that 11 rays travel from the robot to obstacles, similar to sparse LiDAR readings. We train a network that maps raw depth images to predicted ray distances (see also Section VII), and the ray distances serve as part of the observations for the agile policy and the RA value network.

To summarize, as shown in Figure 2 (a), ABS needs to train four modules all in simulation:

- 1) The *agile policy* (Section IV) is trained via RL to reach the goal without collisions. We design goal-reaching rewards to encourage the most agile motor skills.
- 2) The *RA value network* (Section V) is trained to indicate the safety for the agile policy. We use a data-driven method to train it based on the RA bellman equation (Equation (5)), and collect data in simulation by rolling out the agile policy.
- 3) The *recovery policy* (Section VI) is trained to track twist commands rapidly from high-speed movements.
- 4) The *ray-prediction network* (Section VII) is trained to predict ray distance observations from depth images. We collect synthetic depth images and ray distances in simulation by rolling out the agile policy.

All of the four modules are directly deployed in the real world after training.

IV. LEARNING AGILE POLICY

As mentioned in Section III-C, we train an agile policy to achieve high agility amidst obstacles. Previous works on learning agile locomotion typically employ the velocity-tracking formulation [48, 58], *i.e.*, to track velocity commands on open, flat terrains. However, designing a navigation planner for these velocity-tracking policies in cluttered environments can be non-trivial. To ensure safety, the planner may have to be conservative and unable to fully unleash the locomotion policy's agility.

Instead, we use the goal-reaching formulation to maximize the agility, inspired by [54, 82]. Specifically, we train the robot to develop sensorimotor skills that enable it to reach specified goals within the episode time without collisions. The agility is also encouraged by a reward term pursuing high velocity in the base frame. By doing so, the robot naturally learns to achieve maximum agility while avoiding collisions.

This section presents the details of our agile policy learning. A detailed comparison between goal-reaching and velocity-tracking formulations for agility will be presented in Section IX-A1.

A. Observation Space and Action Space

The observation space of the agile policy consists of the foot contacts $c_{f \in \{1,2,3,4\}}$, the base angular velocities ω , the projected gravity in the base frame g , the goal commands G^c (*i.e.*, the relative position and heading of the goal) in the base frame, the time left of the episode $T - t$, the joint positions q , the joint velocities $\dot{q}$, the actions a of the previous frame, and the exteroception (*i.e.*, log values of the ray distances) R . Here we omit the step-based timestamps ($t - 1$ for the actions and t otherwise) for brevity. We refer to the collection of all these variables as o^{Agile} .

Among these observations, only g and G^c require the state estimators for respectively orientation and odometry. All other values are available from raw sensor data without cumulative drifts. The IMU-based orientation estimation for g (*i.e.*, roll and pitch) is usually very accurate, and our policy can effectively handle the odometry drift (as we can even suddenly change the goals in the run, see Section IX-D). Therefore, our agile policy is robust to inaccurate state estimators which can be problematic for model-based controllers [6, 34, 18].

The action space of the agile policy consists of 12-d joint targets. A PD controller tracks these joint targets a by converting them to joint torques:

$$\tau = K_p(a - q) - K_d\dot{q}. \quad (6)$$

A fully-connected MLP maps the observations o^{Agile} to the actions a .

B. Rewards

Our reward function is the summation of multiple terms:

$$r = r_{\text{penalty}} + r_{\text{task}} + r_{\text{regularization}}, \quad (7)$$

where each term can be further divided into subterms as follows.

1) *Penalty Rewards*: We use a simple penalty design:

$$r_{\text{penalty}} = -100 \cdot \mathbb{1}(\text{undesired collision}), \quad (8)$$

where undesired collision refers to collisions on the base, thighs, and calves, and horizontal collisions on the feet.

2) *Task Rewards*: The task rewards are:

$$r_{\text{task}} = 60 \cdot r_{\text{possoft}} + 60 \cdot r_{\text{postight}} + 30 \cdot r_{\text{heading}} - 10 \cdot r_{\text{stand}} + 10 \cdot r_{\text{agile}} - 20 \cdot r_{\text{stall}}, \quad (9)$$

i.e., a soft position tracking term r_{possoft} to encourage the exploration for goal reaching, a tight position tracking term r_{postight} to reinforce the robot to stop at the goal, a heading tracking term r_{heading} to regulate the robot's heading near the goal, a standing term r_{stand} to encourage a standing posture at the goal, an agile term r_{agile} to encourage high velocities, and a stall term r_{stall} to penalize the waiting behaviors. These terms ensure that the robot should reach the goal with appropriate heading and posture as fast as possible while wasting no time.

To be specific, our tracking terms (r_{possoft} , r_{postight} , r_{heading}) are in the same form as shown below, inspired by [83] where RL-based navigation planners are learned:

$$r_{\text{track (possoft/postight/heading)}} = \frac{1}{1 + \left\| \frac{\text{error}}{\sigma} \right\|^2} \cdot \frac{\mathbb{1}(t > T - T_r)}{T_r}, \quad (10)$$

where σ normalizes the tracking errors, T is the episode length, and T_r is a time threshold. By doing so, the robot only needs to reach the goal before $T - T_r$ to maximize the tracking rewards, free from explicit motion constraints such as target velocities that may limit the agility. For the soft position tracking, we have $\sigma_{\text{soft}} = 2$ m and $T_r = 2$ s with the error being the distance to the goal. For the tight position tracking, we have $\sigma_{\text{tight}} = 0.5$ m and $T_r = 1$ s. For the heading tracking, we have $\sigma_{\text{heading}} = 1$ rad and $T_r = 2$ s with the error being the relative yaw angle to the goal heading. We further disable r_{heading} when the distance to the goal is larger than σ_{soft} so that collision avoidance is not affected.

The standing term is defined as

$$r_{\text{stand}} = \|\bar{q} - q\|_1 \cdot \frac{\mathbb{1}(t > T - T_{r,\text{stand}})}{T_{r,\text{stand}}} \cdot \mathbb{1}(d_{\text{goal}} < \sigma_{\text{tight}}), \quad (11)$$

where $\bar{q}$ is the nominal joint positions for standing, $T_{r,\text{stand}} = 1$ s, and d_{goal} is the distance to the goal.

The agile term is the core term that encourages the agile locomotion. It is defined as

$$r_{\text{agile}} = \max \left\{ \text{ReLU} \left(\frac{v_x}{v_{\text{max}}} \right) \cdot \mathbb{1}(\text{correct direction}), \mathbb{1}(d_{\text{goal}} < \sigma_{\text{tight}}) \right\}, \quad (12)$$

where v_x is the forward velocity in the robot base frame, $v_{\text{max}} = 4.5$ m/s is an upper bound of v_x that cannot be

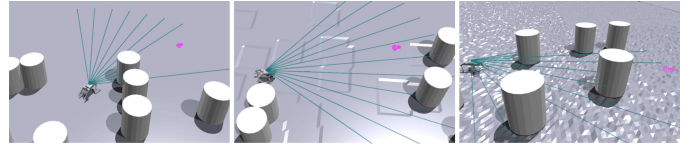


Fig. 3. Example training environments. The magenta points indicate the goals, and the bluegreen lines indicate the exteroceptive ray observations. Terrains from left to right: flat, low stumbling blocks, and rough.

reached (based on the hardware datasheet), and the “correct direction” means that the angle between the robot heading and the robot-goal line is smaller than 105° . To maximize this term, the robot has to either run fast or stay at the goal.

The stall term r_{stall} is 1 if the robot stays static when $d_{\text{goal}} > \sigma_{\text{soft}}$ and the robot is not in the “correct direction”. This term penalizes the robot for time waste.

3) *Regularization Rewards*: The regularization rewards are:

$$\begin{aligned} r_{\text{regularization}} = & -2 \cdot v_z^2 - 0.05 \cdot (\omega_x^2 + \omega_y^2) - 20 \cdot (g_x^2 + g_y^2) \\ & - 0.0005 \cdot \|\tau\|_2^2 - 20 \cdot \sum_{i=1}^{12} \text{ReLU}(|\tau_i| - 0.85 \cdot \tau_{i,\text{lim}}) \\ & - 0.0005 \cdot \|\dot{q}\|_2^2 - 20 \cdot \sum_{i=1}^{12} \text{ReLU}(|\dot{q}_i| - 0.9 \cdot \dot{q}_{i,\text{lim}}) \\ & - 20 \cdot \sum_{i=1}^{12} \text{ReLU}(|q_i| - 0.95 \cdot q_{i,\text{lim}}) \\ & - 2 \times 10^{-7} \cdot \|\ddot{q}\|_2^2 - 4 \times 10^{-6} \cdot \|\dot{a}\|_2^2 - 20 \cdot \mathbb{1}(\text{fly}), \end{aligned} \quad (13)$$

where τ is the joint torques, τ_{lim} is the hardware torque limits, $\dot{q}_{\text{lim}}$ is the hardware joint velocity limits, q_{lim} is the hardware joint position limits, and “fly” refers to when the robot has no contact with the ground. We penalize the “fly” cases as they make the robot base uncontrollable, threatening the system's safety.

C. Training in Simulation

1) *Simulator*: We use the GPU-based Isaac Gym simulator [47] which supports us to train 1280 environments in parallel with the PPO algorithm [56].

2) *Terrains*: We train the agile policy with randomized terrains following a curriculum to facilitate learning. To prevent unstable gaits that over-exploits the simulation dynamics, the terrains are randomly sampled to be flat, rough, or low stumbling blocks, as shown in Figure 3. As the difficulty level goes up from 0 to 9, the rough terrains and the stumbling blocks have larger height difference from 0 cm to 7 cm.

3) *Obstacles*: We train the policy with cylinders of 40 cm radius. For each episode we have 0~8 obstacles randomly distributed in a 11 m $\times$ 5 m rectangle that covers the origin and the goal. To facilitate learning, we also apply a curriculum where higher difficulty levels have more obstacles.

4) *Domain Randomization*: We do domain randomization [67] for sim-to-real transfer. The randomized settings are listed in Table II. Among these few terms, two are critical: the illusion, and the ERFI-50. The illusion makes the policy more robust to unseen geometries such as walls: it overwrites the observed ray distances by random values subject to $\mathcal{U}(d_{\text{goal}} + 0.3, \text{ray distance})$ if they are larger than

TABLE II
DOMAIN RANDOMIZATION SETTINGS FOR AGILE POLICY TRAINING

Term	Value
Observation	
Illusion	Enabled
Joint position noise	$\mathcal{U}(-0.01, 0.01)$ rad
Joint velocity noise	$\mathcal{U}(-1.5, 1.5)$ rad/s
Angular velocity noise	$\mathcal{U}(-0.2, 0.2)$ rad/s
Projected gravity noise	$\mathcal{U}(-0.05, 0.05)$
log(ray distance) noise	$\mathcal{U}(-0.2, 0.2)$
Dynamics	
ERFI-50 [8]	0.78 N m $\times$ difficulty level
Friction factor	$\mathcal{U}(0.4, 1.1)$
Added base mass	$\mathcal{U}(-1.5, 1.5)$ kg
Joint position biases	$\mathcal{U}(-0.08, 0.08)$ rad
Episode	
Episode length	$\mathcal{U}(7.0, 9.0)$ s
Initial robot position	$x = 0, y = 0$
Initial robot yaw	$\mathcal{U}(-\pi, \pi)$ rad
Initial robot twist	$\mathcal{U}(-0.5, 0.5)$ m/s or rad/s
Goal Position	$x_{\text{goal}} \sim \mathcal{U}(1.5, 7.5)$ m
	$y_{\text{goal}} \sim \mathcal{U}(-2.0, 2.0)$ m
Goal Heading	$\arctan 2(y_{\text{goal}}, x_{\text{goal}}) + \mathcal{U}(-0.3, 0.3)$ rad

$d_{\text{goal}} + 0.3$. The ERFI-50 proposed by Campanaro et al. [8] implicitly models the motor sim-to-real gaps with random torque perturbations, and we add a curriculum in our work to avoid impeding the early stage of learning. We also randomly bias the joint positions to model the motor encoders' offset errors.

5) *Curriculum*: As mentioned above, we apply a curriculum where difficulty levels can change the terrains, the obstacle distribution, and the domain randomization. For the assignment of difficulty levels, we follow the design of Zhang et al. [82]: when an episode terminates, the robot gets promoted to a higher level if $d_{\text{goal}} < \sigma_{\text{tight}}$, and gets demoted to a lower level if $d_{\text{goal}} > \sigma_{\text{soft}}$. If the robot gets promoted at the highest level, it will go to a random level, following [55].

V. LEARNING AND USING REACH-AVOID VALUES

Although the agile policy learns certain collision avoidance behaviors via corresponding rewards, it does not ensure safety. To safeguard the robot, we propose to use RA values to predict the failures, and then a recovery policy can save the robot based on the RA values.

Inspired by Hsu et al. [30], we learn RA values in a model-free way, contrasting typical approaches of model-based reachability analysis [2]. This better suits the model-free RL-based policies. Also different from [30], we do not learn the global RA values, but make it policy-conditioned, as mentioned in Section III-B4. The learned RA value function will predict only the agile policy's failures based on the observations.

A. Learning RA Values

To avoid overfitting in high dimensions and make the RA values generalize, we use a reduced set of observations as the inputs of the RA value function:

$$o^{\text{RA}} = [v; \omega; G_{x,y}^c; R], \quad (14)$$

i.e., the base twists, the goal (x, y) position in the robot frame, and the exteroception. These components are centroidal observations that significantly influence safety and goal reaching.

On the other hand, we don't use joint-level observations (such as q and $\dot{q}$) here because they are high-dimensional and less pertinent to goal reaching. We train an RA value network $\hat{V}$ to approximate the RA values:

$$V_{\text{RA}}^{\pi^{\text{Agile}}}(s) \approx \hat{V}(o^{\text{RA}}). \quad (15)$$

Based on Equation (5), we minimize the following loss for each episode with gradient descent:

$$L = \frac{1}{T} \sum_{t=1}^T \left(\hat{V}(o_t^{\text{RA}}) - \hat{V}^{\text{target}} \right)^2, \quad (16)$$

where

$$\begin{aligned} \hat{V}^{\text{target}} = & \gamma_{\text{RA}} \max \left\{ \zeta(s_t), \min \{ l(s_t), \hat{V}^{\text{old}}(o_{t+1}^{\text{RA}}) \} \right\} \\ & + (1 - \gamma_{\text{RA}}) \max \{ l(s_t), \zeta(s_t) \}, \end{aligned} \quad (17)$$

and we set the discount factor $\gamma_{\text{RA}} = 0.999999$ to best approximate $\mathcal{RA}^{\pi}(\Theta; \mathcal{F})$ since $V_{\text{RA}}^{\pi}(s)$ converges to $V_{\text{RA}^*}^{\pi}(s)$ as γ_{RA} approaches 1. $\hat{V}^{\text{old}}$ refers to $\hat{V}$ from previous iteration, and we set $\hat{V}^{\text{old}}(o_{T+1}^{\text{RA}}) = +\infty$.

Differing from [30], our approach learns policy-conditioned reach-avoid values instead of solving policy-agnostic global reach-avoid value of the entire system dynamics which involves another value minimization problem over $\mathcal{A}$. Our method offers several advantages: 1) simplicity: as highlighted in Equation (5), this simplicity arises from avoiding the need to solve for the lowest value of the next state across the entire action space. 2) two-stage offline learning: our approach can be learned in a two-stage offline manner. This involves first collecting policy trajectories and then training the policy-conditioned reach-avoid value. This two-stage process enhances stability compared to the online training method presented in [30].

B. Implementation

According to [30], $l(s)$ and $\zeta(s)$ should be Lipschitz-continuous for theoretical guarantees. In our implementation, we define the $l(s)$ as

$$l(s) = \tanh \log \frac{d_{\text{goal}}}{\sigma_{\text{tight}}}, \quad (18)$$

thereby making it Lipschitz-continuous, bounding it with $(-1, 1)$, and setting $d_{\text{goal}} \leq \sigma_{\text{tight}}$ as "reach".

Regarding failures, we naturally have

$$\zeta(s) = 2 * \mathbb{1}(\text{undesired collision}) - 1. \quad (19)$$

However, this definition violates the Lipschitz continuity. Hence, we soften the function in a hindsight way: when an undesired collision happens, the ζ values for the last 10 timesteps are relabelled to be $-0.8, -0.6, \dots, 0.8, 1.0$.

For RA dataset sampling, we make the obstacles distribute as in the highest difficulty level during the training of the agile policy. We roll out our trained agile policy for 200k episodes, and collect these trajectories for RA learning.

Figure 4 visualizes the learned RA values for a specific set of obstacles. As the robot's velocity changes, the landscape of RA values changes accordingly. The sign of the RA values reasonably indicates the safety for the agile policy.

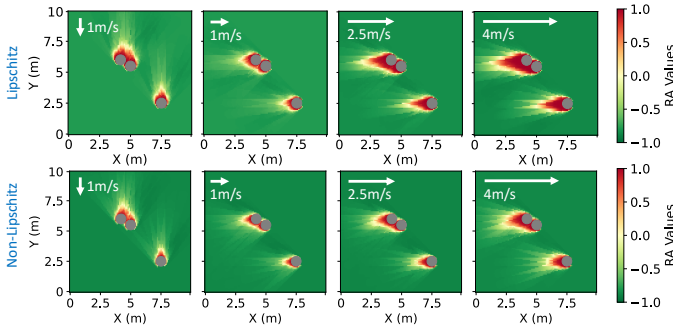


Fig. 4. Visualization of $\hat{V}$ with different linear velocities and 2D positions relative to the 3 fixed obstacles. The angular velocities are set to zero, and the relative goal commands are set to 5 m ahead of the robot. The grey circles represent the obstacles, and the colors represent the values of $\hat{V}$ at corresponding 2D positions. The first row presents the RA values trained with the softened failure function ζ , while the second row uses the raw one in Equation (19). Without softening ζ to approach the Lipschitz continuity, the value estimation fails to indicate collisions on the sides of obstacles and has local minima in front of the obstacles, compromising safety.

C. Using RA Values for Recovery

RA values provide a failure prediction conditioned on the agile policy, and we propose to use RA values to guide the recovery policy. To be specific, the robot decides the optimal twist to avoid collisions using the RA value function, and employs the recovery policy to track these twist commands. The recovery policy is triggered as a back-up shielding policy if and only if $\hat{V}(o^{\text{RA}}) \geq V_{\text{threshold}}$. We set $V_{\text{threshold}} = -0.05$ to compensate for learning errors without causing over-conservative shielding.

During recovery, we assume that the recovery policy is well-trained so that the robot twist is close to the command

$$tw^c = [v_x^c, v_y^c, 0, 0, 0, \omega_z^c], \quad (20)$$

and the robot should try to get closer to the goal if its twist is safe given the goal and the exteroception. Therefore, the twist command is obtained from the optimization:

$$tw^c = \arg \min d_{\text{goal}}^{\text{future}} \text{ s.t. } \hat{V}([tw^c; G_{x,y}^c; R]) < V_{\text{threshold}}, \quad (21)$$

and $d_{\text{goal}}^{\text{future}}$ refers to the approximate distance to the goal after tracking the twist command for a small amount of time $\delta t = 0.05$ s. This is calculated based on the linearized integral of the robot displacement in the base frame:

$$\begin{aligned} \delta x &= v_x^c \delta t - 0.5 v_y^c \omega_z^c \delta t^2, \\ \delta y &= v_y^c \delta t + 0.5 v_x^c \omega_z^c \delta t^2. \end{aligned} \quad (22)$$

In our practice, gradient descent with a Lagrangian multiplier on the constraint can solve Equation (21) within 5 steps when initialized with the current twist, thereby enabling real-time deployment. A visualization of the twist optimization process is given in Figure 8 where the searched twist consistently satisfies the safety constraint (i.e., $\hat{V} < V_{\text{threshold}}$).

VI. LEARNING RECOVERY POLICY

The recovery policy is intended to make the robot track a given twist command as fast as possible so that it can function as a backup shielding policy, as mentioned in Section V.

A. Observation Space and Action Space

The observation space of the recovery policy differs from the agile policy in that it tracks twist commands and it does not need exteroception. The recovery policy's observation o^{Rec} consists of: the foot contacts c_f , the base angular velocities ω , the projected gravity in the base frame g , the twist commands tw^c (only non-zero variables), the joint positions q , the joint velocities $\dot{q}$, and the actions a of the previous frame.

The action space of the recovery policy is exactly the same as that of the agile policy: the 12-d joint targets. We also use an MLP as the policy network.

B. Rewards

Similar to the agile policy, the reward functions for the recovery policy also consist of three parts: the penalty rewards, the task rewards, and the regularization rewards. The regularization rewards and the penalty rewards remain the same, except that we allow knee contacts with the ground for maximum deceleration (e.g., Figure 1 (a)).

The task rewards are for twist tracking:

$$r_{\text{task}} = 10 \cdot r_{\text{linvel}} - 0.5 \cdot r_{\text{angvel}} + 5 \cdot r_{\text{alive}} - 0.1 \cdot r_{\text{posture}}, \quad (23)$$

i.e., a term for tracking v_x^c and v_y^c , a term for tracking ω_z^c , a term for staying alive, and a term for maintaining a posture to seamlessly switch back to the agile policy.

To be specific, we have

$$r_{\text{linvel}} = \exp \left[-\frac{(v_x - v_x^c)^2 + (v_y - v_y^c)^2}{\sigma_{\text{linvel}}^2} \right], \quad (24)$$

where we set $\sigma_{\text{linvel}} = 0.5$ m/s. For the angular velocity,

$$r_{\text{angvel}} = \|\omega_z - \omega_z^c\|_2^2, \quad (25)$$

which provides a softer landscape near the command than r_{linvel} . The alive term is simply

$$r_{\text{alive}} = 1 \cdot \mathbb{1}(\text{alive}). \quad (26)$$

The posture term is

$$r_{\text{posture}} = \|q - \bar{q}_{\text{rec}}\|_1, \quad (27)$$

where $\bar{q}_{\text{rec}}$ is a nominal standing pose with low height allowing the robot to switch back to the agile policy seamlessly.

C. Training in Simulation

The simulation settings for training the recovery policy are similar to those for the agile policy. The differences lie in:

1) *Domain Randomization*: The observation noises and the dynamic randomization do not change. The episode length is changed to 2 s, and there are randomized initial roll and pitch angles subject to $\mathcal{U}(-\pi/6, \pi/6)$ rad. The randomization ranges are also changed for initial $v_x \sim \mathcal{U}(-0.5, 5.5)$ m/s and initial $\omega \sim \mathcal{U}(-1.0, 1.0)$ rad/s. These changes better accommodate the states that can trigger the recovery policy during the agile running. The ranges of sampling commands are $v_x^c \sim \mathcal{U}(-1.5, 1.5)$ m/s, $v_y^c \sim \mathcal{U}(-0.3, 0.3)$ m/s, and $\omega_z^c \sim \mathcal{U}(-3.0, 3.0)$ rad/s.

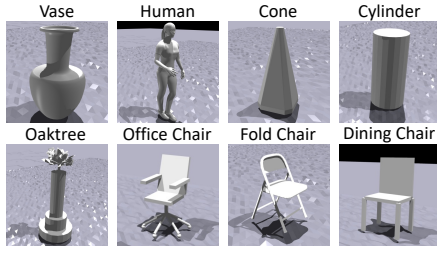


Fig. 5. Various obstacles used for ray-prediction data collection.

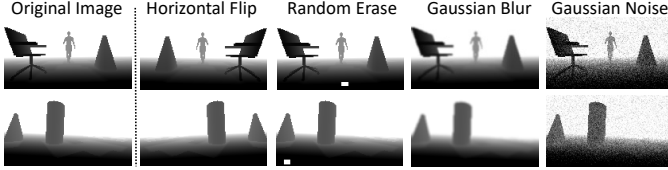


Fig. 6. Illustration of four kinds of image augmentation used for depth-based ray-prediction training.

2) *Curriculum*: The curriculum still exists for terrains and domain randomization. However, the assignment is changed: the robot gets promoted if the velocity tracking error is smaller than $0.7\sigma_{\text{linvel}}$, and gets demoted if it falls over.

VII. PERCEPTION

As mentioned in Section IV and Section V, both the agile policy and the RA value network use the exteroceptive 11-d ray distances as part of the observations, with access to their ground truth values during training. These rays are horizontally cast from the robot base, with directions evenly spaced in $[-\frac{\pi}{4}, \frac{\pi}{4}]$.

However, such ray distances are not directly available during deployment, and we need to train a ray-prediction network to predict them from depth images, as mentioned in Section III-C. Such a design leads to the following benefits:

- 1) We only need to tune the ray-prediction network to handle high-dimensional image noises by data augmentation.
- 2) The representation is highly interpretable, allowing humans to supervise.
- 3) The agile policy and the RA value network are easier to train with low-dimensional inputs.
- 4) Compared to costly image rendering in simulation, the ray distances are easy to compute and save training time.

Besides, although ray distances are similar to sparse LiDAR readings, we use cameras instead of LiDARs because a lightweight low-cost camera can easily reach a high FPS, which is important in high-speed collision avoidance.

We present details for training the ray-prediction network in this section.

A. Data Collection

To train our ray-prediction network, we collect a dataset of pairs of depth images and ray distances (as shown in Figure 2 (a)) by running the agile policy in simulation. The ray-prediction network can then be trained in a supervised

way. To facilitate generalization, as shown in Figure 5, we replace the cylinders with objects of different shapes during data collection.

B. Data Augmentation for Sim-to-Real Transfer

The real-world depth images collected from cameras are far more noisy than the rendered depth images in simulation [28]. To make the ray-prediction network adapt better to real-world depth images, we apply four data augmentation techniques during training, as shown in Figure 6: 1) horizontal flip; 2) random erase; 3) Gaussian blur; 4) Gaussian noise. For deployment, we apply hole filling [62] to further reduce the gap of depth images between the simulation and the real world.

C. Other Implementation Details

To make the network focus more on close obstacles, we take the logarithm of depth values as the NN inputs, and the logarithm of ray distances as the outputs, with the mean squared error as the loss function.

We finetune ResNet-18 [27] with pretrained weights to train the model. The images are downsampled to $[160, 90]$ resolution both in simulation and during deployment.

VIII. EXPERIMENTS

A. Baselines

For experimental results, we consider three settings:

- 1) Our ABS system, with both the agile policy and the recovery policy;
- 2) Our agile policy π^{Agile} only;
- 3) “LAG”: we use PPO-Lagrangian [53] to train end-to-end safe RL policies with the agile policy’s formulation.

By comparing (2) and (3), we can see how agility and safety trade off without external modules, forming a boundary of agility and safety. With the help of RA values and the recovery policy, we expect (1) to break this boundary with a high safety gain: it should be as agile as (2) in safe cases, and shield the robot in risky cases.

Note that the three settings here are all based on what we propose. A detailed comparison between our agile policy and the previous state-of-the-art (SOTA) agile running policy [48] is made in Section IX-A1.

B. Simulation Experiments

1) *Quantitative results*: We test the policies trained with different settings in simulation. To better show the agility-safety boundary, we introduce 3 variants for each setting: an aggressive one (“-a”) doubling the agile reward term r_{agile} , a nominal one (“-n”), and a conservative one (“-c”) halving the r_{agile} . Regarding the obstacles, we distribute eight obstacles within a $5.5 \text{ m} \times 4 \text{ m}$ rectangle (during training it was $11 \text{ m} \times 5 \text{ m}$), so the test cases are in distribution but much harder than most cases during training.

The results are reported in Table III and Figure 7. There are three possible outcomes for an episode: success, collision, or timeout. Trajectories that do not trigger success or collision criteria within the episode length are labelled as “timeout”.

TABLE III
BENCHMARKED COMPARISON IN SIMULATION

	Success Rate (%)	Collision Rate (%)	Timeout Rate (%)	$\bar{v}_{\text{peak}}$ of Success (m/s)	$\bar{v}$ of Success (m/s)
ABS-a	78.9 $\pm$ 1.4	4.4 $\pm$ 0.5	16.7 $\pm$ 1.9	3.74 $\pm$ 0.02	2.15 $\pm$ 0.04
ABS-n	79.1 $\pm$ 4.4	5.7 $\pm$ 2.9	15.2 $\pm$ 2.1	3.48 $\pm$ 0.06	2.08 $\pm$ 0.01
ABS-c	85.8$\pm$5.6	2.9$\pm$0.7	11.3 $\pm$ 5.1	2.98 $\pm$ 0.12	1.87 $\pm$ 0.03
π^{Agile} -a	73.3 $\pm$ 4.3	26.1 $\pm$ 4.4	0.6$\pm$0.1	3.83$\pm$0.03	2.55$\pm$0.03
π^{Agile} -n	77.3 $\pm$ 4.2	21.7 $\pm$ 3.9	1.0 $\pm$ 0.4	3.55 $\pm$ 0.04	2.39 $\pm$ 0.04
π^{Agile} -c	83.2$\pm$1.7	15.5 $\pm$ 2.0	1.3 $\pm$ 0.6	3.04 $\pm$ 0.13	2.04 $\pm$ 0.08
LAG-a	82.5$\pm$6.0	10.9 $\pm$ 2.6	6.6 $\pm$ 4.5	2.70 $\pm$ 0.13	1.69 $\pm$ 0.09
LAG-n	77.4 $\pm$ 11.5	9.1 $\pm$ 1.8	13.5 $\pm$ 13.0	2.45 $\pm$ 0.07	1.41 $\pm$ 0.03
LAG-c	49.1 $\pm$ 8.4	7.4 $\pm$ 2.7	43.5 $\pm$ 11.1	2.45 $\pm$ 0.10	1.12 $\pm$ 0.08

*Bold values: the mean falls within the range of top1's mean $\pm$ top1's std.

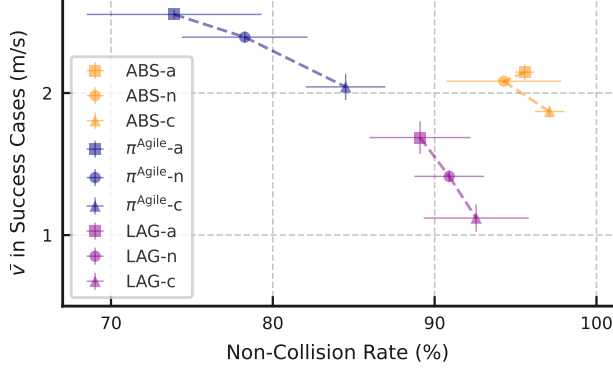


Fig. 7. Illustration of agility-safety trade-off in benchmarked comparison. Agility is quantified by the average speed achieved in success cases while safety is represented by the non-collision rate. Points indicate the mean values, and error bars indicate the std values.

We report the success rate, the collision rate, the timeout rate, the average peak velocity for success cases, and the average speed for success cases as the metrics. For each setting, the mean and std values are calculated over 3 policies trained with different seeds, and the metrics are obtained via testing for 10k random episodes.

The results indicate that, no matter how the reward weights are tuned or whether the RL algorithm is constrained by safe exploration, the agility and the safety trade off within a boundary. Yet, with our RA values and the recovery policy as a safeguard, we can break this boundary and get a substantial improvement in safety at the cost of only a minor decrease in agility.

Note that the variants are only introduced to show the boundary here. In the following parts, we will only use the nominal ones.

2) *Example Case:* We present an example case of ABS and other baselines in simulation, where the robot starting from (0, 0) needs to run through 8 obstacles to reach the goal (7, 0), as shown in Figure 8. The robot needs to first go through an open space, followed by two tight spaces, and then another open space. In this case, the π^{Agile} baseline runs fast but collides near the second tight space. The LAG baseline runs much slower than ABS. In contrast, our proposed ABS runs fast in the open spaces, and slows down in the tight spaces for safety thanks to the shielding of RA values and the recovery policy. Figure 8 (c) demonstrates the RA value landscape with respect to twist commands when the recovery policy is

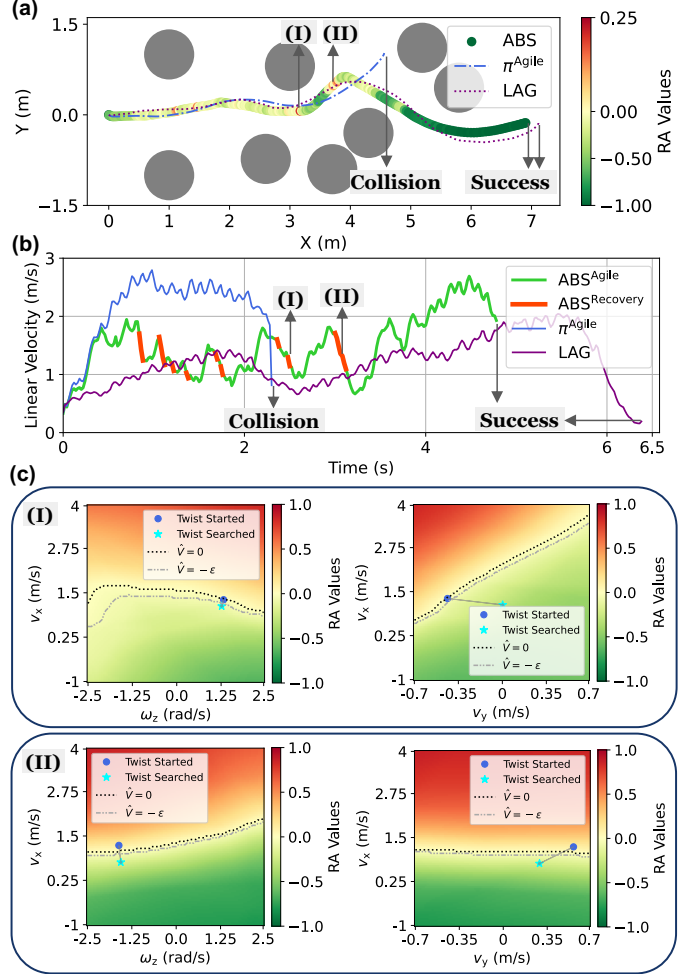


Fig. 8. An example case in simulation where π^{Agile} fails to reach the goal. a) Trajectories of ABS and other baselines, with RA values visualized for ABS. b) The velocity-time curves showing that ABS is much faster than the LAG baseline. c) Illustrations of the RA value landscape when the recovery policy is triggered at (I) and (II), projected in the $v_x - \omega_z$ plane and the $v_x - v_y$ plane. We show the initial twist before search (i.e., the current twist of the robot base) and the searched commands based on Equation (21).

activated, where the searched twist consistently satisfies the safety constraint (i.e., $\hat{V} < V_{\text{threshold}}$).

C. Real-World Experiments

1) *Hardware setup:* We use the Unitree Go1 for our experiments. The robot is equipped a Jetson Orin NX for onboard

	Success	Collision	Time Cost
Indoor (a)			
ABS	9/10	1/10	5.91 s
ABS (only π_{Agile})	7/10	3/10	5.06 s
LAG	8/10	2/10	6.80 s
Indoor (b)			
ABS	10/10	0/10	4.75 s
ABS (only π_{Agile})	7/10	3/10	3.74 s
LAG	9/10	1/10	6.13 s
Outdoor			
ABS	10/10	0/10	4.46 s
ABS (only π_{Agile})	9/10	1/10	4.15 s
LAG	9/10	1/10	6.05 s
Speed Test			
	ABS	LAG	
Peak Speed	3.1 m/s	2.1 m/s	

Fig. 9. We evaluate ABS and other baselines in the real world with two indoor testbeds, one outdoor testbed, and repetitive speed tests. Indoor (a) is a dim and narrow corridor, Indoor (b) is a hall with furnitures, and Outdoor is an open space on the playground with few obstacles. ABS achieves the best safety across three testbeds, with faster speeds compared to the LAG baseline.

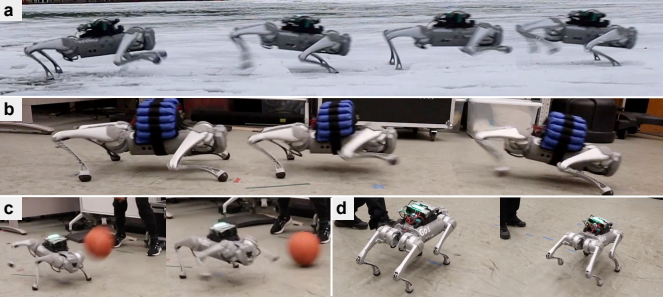


Fig. 10. Robustness Tests of our ABS system, a) in snowy terrain, b) bearing a 12-kg payload, c) against a ball hit when running, and d) withstand a kick when standing at the goal.

computation and a ZED Mini Stereo Camera for depth and odometry sensing. We employ the ZED odometry module to online update the relative goal commands for the agile policy, with its difference as the velocity estimation. We use Unitree’s built-in PD controller, with $K_p = 30$ and $K_d = 0.65$.

2) *Results*: Across two indoor and one outdoor testbeds, ABS demonstrates superior overall performance as shown in Figure 9, achieving the highest success rates and the lowest collision rates. Specifically, ABS consistently scores either 9 or 10 out of 10 in success rates across all environments, with minimal collisions, indicating robustness and reliability in the real world.

Without the safety shield, the agile policy π_{Agile} achieves

TABLE IV
GOAL-REACHING POLICY V.S. VELOCITY-TRACKING POLICY

Term	Our π_{Agile}	Rapid [48]
Gait pattern	gallop	near trot
Max #. uncontrollable DoFs	1	3
Peak vel. in simulation	4.0 m/s	4.1 m/s
Peak torque in simulation	23.5 Nm	35.5 Nm
Peak joint vel. in simulation	22.0 rad/s	30.0 rad/s
Peak vel. in real world	3.1 m/s	2.5 m/s
Collision avoidance	as trained	need high-level commands
Fully unleashed agility	as trained	non-trivial for high level
Changing vel. for steering	in distribution	out of distribution
Curriculum learning	straightforward	carefully designed

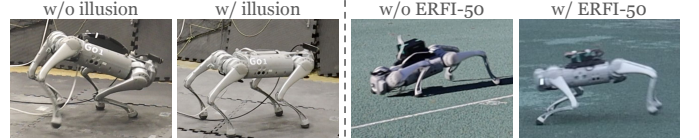


Fig. 11. Effects of illusion and ERFI-50 randomization. The robot will tremble near a wall without illusion randomization and will hit the ground during running without ERFI-50 randomization.

the fastest running speed at the cost of more collisions. LAG outperforms π_{Agile} in safety but has slower speeds, and falls short in both safety and agility compared to ABS. ABS achieves high speed with high safety, and generalizes to dynamic obstacles, as shown in Figure 1.

3) *Robustness*: Our ABS system can work on the slippery icy snow, bear a 12-kg payload (equal to its own weight), and withstand perturbations, as shown in Figure 10. These tests demonstrate the robustness of our system.

IX. EXTENSIVE STUDIES AND ANALYSES

A. Maximizing Agility

1) *Goal-Reaching v.s. Velocity-Tracking*: Velocity-tracking is the most commonly used formulation of locomotion controllers [55, 48, 50, 41], and is also adopted for our recovery policy. However, for the agile policy, we claim that goal-reaching is a better choice because it does not decouple locomotion and navigation for collision avoidance and can fully unleash the agility that is learned. Moreover, we empirically find that the goal-reaching formulation benefits sim-to-real transfer as it finds a better gait pattern for high-speed running.

With the SOTA agile velocity-tracking policy in [48] as a baseline (referred to as “rapid”), we make detailed comparisons in Table IV. For fair comparison, we train “rapid” with the combination of our regularization rewards and the task rewards in [48], use the same action space for two policies, and remove the temporal information and system identification in [48].

2) *Effects of illusion and ERFI-50 randomization*: Two key components we add in domain randomization to help sim-to-real transfer is the illusion and the ERFI-50. As shown in Figure 11, without the illusion, the robot will sometimes tremble near a wall which it has never seen in simulation. Without ERFI-50, the robot will hit the ground with its head during running due to the sim-to-real gap in motor dynamics.

TABLE V
EFFECTS OF DIFFERENT $V_{\text{THRESHOLD}}$ ON ABS

$V_{\text{threshold}}$	-0.001	-0.01	-0.05	-0.1
Success Rate (%)	78.0±2.1	78.1±3.4	79.1±4.4	75.8± 2.0
Collision Rate (%)	5.0±0.6	5.8±1.9	5.7±2.9	4.3±0.6
$\bar{v}_{\text{peak}}$ of Success (m/s)	3.42±0.06	3.46±0.08	3.48±0.06	3.42±0.05
$\bar{v}$ of Success (m/s)	2.08±0.02	2.08±0.01	2.08±0.01	2.05± 0.03

TABLE VI
EFFECTS OF SOFTENED FAILURE INDICATOR ON ABS

	ABS w/ softened ζ	ABS w/o softened ζ	π^{Agile}
Success Rate (%)	79.1±4.4	81.7±1.3	77.3±4.2
Collision Rate (%)	5.7±2.9	14.7±1.5	21.7±3.9
$\bar{v}_{\text{peak}}$ of Success (m/s)	3.48±0.06	3.45±0.06	3.55±0.04
$\bar{v}$ of Success (m/s)	2.08±0.01	2.27±0.03	2.39±0.04

B. Extensive Studies on RA Values

1) *Selecting safety threshold*: For safety shielding, we choose $V_{\text{threshold}} = -0.05$. Theoretically, this threshold captures the conservativeness of the switching strategy and the discrepancy between the true reach-avoid value and its learned approximation. Our framework, however, demonstrates robustness to the selection of $V_{\text{threshold}}$, as evidenced by Table VI where scanning $V_{\text{threshold}}$ from -0.001 to -0.1 brings no significant change in the overall performance. For a big value (0.1 is considered big given that V is bounded between -1 and 1), the collision rate slightly decreases as expected whereas the success rate also slightly decreases.

2) *Soft Lipschitz continuity for the failure indicator*: In [20], the Lipschitz continuity is used to prove the existence and uniqueness of a solution to the value function. In our paper, we soften the discrete collision indicator to approach the Lipschitz continuity, and ablate its effects here. As shown in Table V, this technique significantly enhances the safety of our system while slightly increasing the conservativeness, corroborating the observations made in Figure 4.

3) *Can RA values shield LAG baseline*: Given that our framework is general in shielding the goal-reaching policy with RA values and the recovery policy, it can also shield the LAG baseline, which makes a variant of ABS with a cost-critic. We present the results of shielded LAG in Table VII, showcasing improved safety over LAG despite slightly reduced agility. The LAG+RA setting can still reach an average speed of > 1 m/s, demonstrating the power of our ABS framework compared to existing works.

C. Enhancing Perception Training

In refining our ray-prediction network training, we systematically examine several factors: 1) network architecture, 2) pretrained weights, and 3) data augmentation. The comparative results, detailed in Table VIII, underscore the significance of both pretrained weights and data augmentation in enhancing the accuracy of the network.

Regarding the network size, larger networks improve the prediction accuracy at the cost of inference time. Note that the reported inference time in Table VIII was measured on Jetson Orin NX exclusively for perception inference. In practical deployment where computational resources are shared among

TABLE VII
EFFECTS OF RA SHIELDING ON LAG

	ABS	LAG	LAG+RA
Success Rate (%)	79.1±4.4	77.4±11.5	70.5±7.7
Collision Rate (%)	5.7±2.9	9.1±1.8	2.8±1.2
$\bar{v}_{\text{peak}}$ of Success (m/s)	3.48±0.06	2.45±0.07	2.40±0.11
$\bar{v}$ of Success (m/s)	2.08±0.01	1.41±0.03	1.22±0.08

TABLE VIII
PERFORMANCE METRICS FOR DIFFERENT NETWORK ARCHITECTURES AND TRAINING APPROACHES

Architecture	Test Set MSE	Inference Time (ms)
EfficientNet-B0*	3.627×10^{-2}	19
MobileNet-V2*	3.387×10^{-2}	15
ResNet-34	3.081×10^{-2}	14
ResNet-18	3.238×10^{-2}	9
ResNet-18 (w/o pretraining)	3.526×10^{-2}	9
ResNet-18 (w/o augmentation)	3.393×10^{-2}	9

* We use the PyTorch-ONNX pipeline where the implementations of these network architectures may be not fully optimized.

various tasks, the actual update frequency can be considerably lower. For real-time high-speed locomotion, we opt for ResNet-18, balancing accuracy and responsiveness in dynamic environments.

D. Instant Steering via Commands

As mentioned in Section IV, we can change our goal commands even in the run time. Thanks to our single-frame observations and randomization settings, we can easily overwrite goal commands to achieve instant agile steering, as presented in Table IX. This enables direct human involvement similar to the velocity-tracking formulation, and Figure 12 showcases such operations in the real world.

E. Failure Cases and Limitations

First, when the obstacles are too dense and form a local minimum, our policy can easily fail, which is also reflected by the high timeout rates in the results. This is common in local navigation planners [83, 49] though, and a potential solution can be to add memory [74] or introduce a global hint [70].

Second, our generalization to dynamic environments are due to the shielding of RA values and the recovery policy. The RA values are learned with static obstacles, and can only generalize to quasi-static environments. If a dynamic object moves faster than the velocity limit of the recovery policy, the collision may happen. A potential solution is to predict the motions of the obstacles in the future [40, 59].

Third, we limit the robot behaviors to only 2D locomotion and constrain the motions to have no flying phase. For 3D terrains such as stairs and gaps, the problem can be far more challenging because the locomotion skills and the collision avoidance are coupled.

Fourth, implicit system identification techniques [41, 50, 48, 39] can leverage temporal information to represent real-world dynamics and facilitate sim-to-real transfer, but it is non-trivial to incorporate them into our system. This requires a latent embedding of the temporal information, which is hard to deal

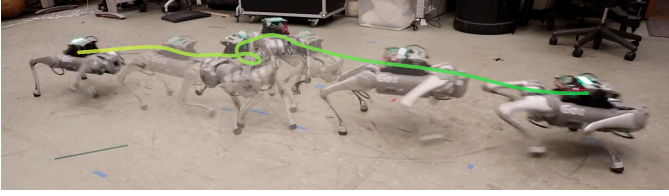


Fig. 12. Steering the robot with a command sequence “Forward”-“Rapid Right Turn”-“Forward”. The robot can reach > 3 m/s when running forward and > 6 rad/s when turning rapidly.

TABLE IX
GOAL COMMANDS FOR INSTANT STEERING

Steering	Goal x (m)	Goal y (m)	Goal Heading (rad)
Forward	5	0	0
Stop	0	0	0
Left Turn	2	1.5	$\frac{\pi}{2}$
Rapid Left Turn	-2	0	$\frac{\pi}{2}$
Right Turn	2	-1.5	$-\frac{\pi}{2}$
Rapid Right Turn	-2	0	$-\frac{\pi}{2}$

with for the RA module. The policy switch can also make the embedding out of distribution for the policies.

Fifth, the vision system needs further improvement. In the Indoor (a) testbed, the only collision of ABS is due to the “undetected” objects by the ray-prediction network as the corridor is quite dim. Beside the network, the system can also be completed by adding more cameras around the body. In this way, the robot may also dodge the obstacles come from behind or the side. Event cameras may also help in highly dynamic scenarios, e.g., when dodging a high-speed ball [19]

X. CONCLUSION

In this paper, we achieve safe high-speed quadrupedal locomotion in cluttered environments. Our framework ABS employs a dual-policy setup where the agile policy enables the robot to run fast, and the recovery policy safeguards the robot. The learned reach-avoid values govern the policy switch and guide the recovery policy. A ray-prediction network provides exteroception representation for the agile policy and the RA value network. Some key takeaways are:

- 1) **Advanced agility:** Decoupled locomotion and navigation constrains the agility of collision-free locomotion in existing works. In contrast, we train an end-to-end agile policy with the local navigation formulation to fully unleash the agility while avoiding obstacles.
- 2) **Safeguarded agility:** The agile policy does not guarantee safety, while using the Lagrangian-based RL or adjusting reward weights only trades off agility and safety. In this paper, we use external shielding modules to help break the trade-off boundary.
- 3) **Model-free safety:** Model-based methods explicitly enforce safety through constraints but struggle with complex dynamics and high-dimensional states. In this paper, we learn RA values in a model-free way and condition them on the policy to simplify learning and enable offline learning from policy rollouts.
- 4) **Guided recovery:** Designing a recovery policy capable of protecting the robot from potential collisions across

various obstacle distributions and states is non-trivial. In this paper, we guide the recovery policy with RA values and their gradients to ensure safety.

- 5) **Perception to motion:** Visual inputs are high-dimensional and noisy, while efficient learning of robust controllers is desired. To this end, we use a low-dimensional exteroception representation to facilitate policy learning and generalization, and train a ray-prediction network to map raw depth images to the representation.

ACKNOWLEDGMENTS

We appreciate Wennie Tabib for supporting hardware experiments, and thank Yuxiang Yang, Yiyu Chen, Yikai Wang and Xialin He for their advice on hardware debugging, and Arthur Allshire for help on software debugging. Special thanks to Andrea Bajcsy and Ziqiao Ma for their assistance in graphics design. This work was in part supported by NSF under grant No. 2144489.

REFERENCES

- [1] Fernando Acero, Kai Yuan, and Zhibin Li. Learning perceptual locomotion on uneven terrains using sparse visual observations. *IEEE Robotics and Automation Letters*, 7(4):8611–8618, 2022.
- [2] Somil Bansal and Claire J Tomlin. Deepreach: A deep learning approach to high-dimensional reachability. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 1817–1824. IEEE, 2021.
- [3] Somil Bansal, Mo Chen, Sylvia Herbert, and Claire J Tomlin. Hamilton-jacobi reachability: A brief overview and recent advances. In *2017 IEEE 56th Annual Conference on Decision and Control (CDC)*, pages 2242–2253. IEEE, 2017.
- [4] Shalabh Bhatnagar and K Lakshmanan. An online actor-critic algorithm with function approximation for constrained markov decision processes. *Journal of Optimization Theory and Applications*, 153:688–708, 2012.
- [5] Gerardo Blede, Matthew J Powell, Benjamin Katz, Jared Di Carlo, Patrick M Wensing, and Sangbae Kim. Mit cheetah 3: Design and control of a robust, dynamic quadruped robot. In *2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2245–2252. IEEE, 2018.
- [6] Michael Bloesch, Christian Gehring, Péter Fankhauser, Marco Hutter, Mark A Hoepflinger, and Roland Siegwart. State estimation for legged robots on unstable and slippery terrain. In *2013 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pages 6058–6064. IEEE, 2013.
- [7] Russell Buchanan, Lorenz Wellhausen, Marko Bjelonic, Tirthankar Bandyopadhyay, Navinda Kottege, and Marco Hutter. Perceptive whole-body planning for multilegged robots in confined spaces. *Journal of Field Robotics*, 38(1):68–84, 2021.

- [8] Luigi Campanaro, Siddhant Gangapurwala, Wolfgang Merkt, and Ioannis Havoutis. Learning and deploying robust locomotion policies with minimal dynamics randomization. *arXiv preprint arXiv:2209.12878*, 2022.
- [9] Mo Chen, Sylvia Herbert, and Claire J Tomlin. Fast reachable set approximations via state decoupling disturbances. In *2016 IEEE 55th Conference on Decision and Control (CDC)*, pages 191–196. IEEE, 2016.
- [10] Richard Cheng, Gábor Orosz, Richard M Murray, and Joel W Burdick. End-to-end safe reinforcement learning through barrier functions for safety-critical continuous control tasks. In *Proceedings of the AAAI conference on artificial intelligence*, volume 33, pages 3387–3395, 2019.
- [11] Xuxin Cheng, Kexin Shi, Ananye Agarwal, and Deepak Pathak. Extreme parkour with legged robots. *arXiv preprint arXiv:2309.14341*, 2023.
- [12] Jia-Ruei Chiu, Jean-Pierre Sleiman, Mayank Mittal, Farbod Farshidian, and Marco Hutter. A collision-free mpc for whole-body dynamic locomotion and manipulation. In *2022 International Conference on Robotics and Automation (ICRA)*, pages 4686–4693. IEEE, 2022.
- [13] Gal Dalal, Krishnamurthy Dvijotham, Matej Vecerik, Todd Hester, Cosmin Paduraru, and Yuval Tassa. Safe exploration in continuous action spaces. *arXiv preprint arXiv:1801.08757*, 2018.
- [14] Jared Di Carlo, Patrick M Wensing, Benjamin Katz, Gerardo Bledt, and Sangbae Kim. Dynamic locomotion in the mit cheetah 3 through convex model-predictive control. In *2018 IEEE/RSJ international conference on intelligent robots and systems (IROS)*, pages 1–9. IEEE, 2018.
- [15] Yanran Ding, Abhishek Pandala, and Hae-Won Park. Real-time model predictive control for versatile dynamic motions in quadrupedal robots. In *2019 International Conference on Robotics and Automation (ICRA)*, pages 8484–8490. IEEE, 2019.
- [16] Helei Duan, Bikram Pandit, Mohitvishnu S Gadde, Bart Jaap van Marum, Jeremy Dao, Chanh Kim, and Alan Fern. Learning vision-based bipedal locomotion for challenging terrain. *arXiv preprint arXiv:2309.14594*, 2023.
- [17] Thomas Dudzik, Matthew Chignoli, Gerardo Bledt, Bryan Lim, Adam Miller, Donghyun Kim, and Sangbae Kim. Robust autonomous navigation of a small-scale quadruped robot in real-world environments. In *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 3664–3671. IEEE, 2020.
- [18] Shamel Fahmi, Geoff Fink, and Claudio Semini. On state estimation for legged locomotion over soft terrain. *IEEE Sensors Letters*, 5(1):1–4, 2021.
- [19] Davide Falanga, Kevin Kleber, and Davide Scaramuzza. Dynamic obstacle avoidance for quadrotors with event cameras. *Science Robotics*, 5(40):eaaz9712, 2020.
- [20] Jaime F Fisac, Mo Chen, Claire J Tomlin, and S Shankar Sastry. Reach-avoid problems with time-varying dynamics, targets and constraints. In *Proceedings of the 18th international conference on hybrid systems: computation and control*, pages 11–20, 2015.
- [21] Jaime F Fisac, Neil F Lugovoy, Vicenç Rubies-Royo, Shromona Ghosh, and Claire J Tomlin. Bridging hamilton-jacobi safety analysis and reinforcement learning. In *2019 International Conference on Robotics and Automation (ICRA)*, pages 8550–8556. IEEE, 2019.
- [22] Kunihiro Fukushima. Visual feature extraction by a multilayered network of analog threshold elements. *IEEE Transactions on Systems Science and Cybernetics*, 5(4):322–333, 1969.
- [23] Magnus Gaertner, Marko Bjelonic, Farbod Farshidian, and Marco Hutter. Collision-free mpc for legged robots in static and dynamic scenes. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 8266–8272. IEEE, 2021.
- [24] Siddhant Gangapurwala, Mathieu Geisert, Romeo Orsolino, Maurice Fallon, and Ioannis Havoutis. Rloc: Terrain-aware legged locomotion using reinforcement learning and optimal control. *IEEE Transactions on Robotics*, 38(5):2908–2927, 2022.
- [25] Ruben Grandia, Farbod Farshidian, Alexey Dosovitskiy, René Ranftl, and Marco Hutter. Frequency-aware model predictive control. *IEEE Robotics and Automation Letters*, 4(2):1517–1524, 2019.
- [26] Ruben Grandia, Fabian Jenelten, Shaohui Yang, Farbod Farshidian, and Marco Hutter. Perceptive locomotion through nonlinear model-predictive control. *IEEE Transactions on Robotics*, 2023.
- [27] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [28] David Hoeller, Lorenz Wellhausen, Farbod Farshidian, and Marco Hutter. Learning a state representation and navigation in cluttered and dynamic environments. *IEEE Robotics and Automation Letters*, 6(3):5081–5088, 2021.
- [29] David Hoeller, Nikita Rudin, Dhionis Sako, and Marco Hutter. Anymal parkour: Learning agile navigation for quadrupedal robots. *arXiv preprint arXiv:2306.14874*, 2023.
- [30] Kai-Chieh Hsu, Vicenç Rúbies Royo, Claire J. Tomlin, and Jaime F. Fisac. Safety and liveness guarantees through reach-avoid reinforcement learning. In *Robotics: Science and Systems XVII*, 2021.
- [31] Kai-Chieh Hsu, Allen Z Ren, Duy P Nguyen, Anirudha Majumdar, and Jaime F Fisac. Sim-to-lab-to-real: Safe reinforcement learning with shielding and generalization guarantees. *Artificial Intelligence*, 314:103811, 2023.
- [32] Jemin Hwangbo, Joonho Lee, Alexey Dosovitskiy, Dario Bellicoso, Vassilios Tsounis, Vladlen Koltun, and Marco Hutter. Learning agile and dynamic motor skills for legged robots. *Science Robotics*, 4(26):eaau5872, 2019.
- [33] Intel. Embree Ray Tracing Library. <https://www.embree.org/>.

- [34] Fabian Jenelten, Jemin Hwangbo, Fabian Tresoldi, C Dario Bellicoso, and Marco Hutter. Dynamic locomotion on slippery ground. *IEEE Robotics and Automation Letters*, 4(4):4170–4176, 2019.
- [35] Fabian Jenelten, Ruben Grandia, Farbod Farshidian, and Marco Hutter. Tamols: Terrain-aware motion optimization for legged systems. *IEEE Transactions on Robotics*, 38(6):3395–3413, 2022.
- [36] Fabian Jenelten, Junzhe He, Farbod Farshidian, and Marco Hutter. Dtc: Deep tracking control. *Science Robotics*, 9(86):eadh5401, 2024.
- [37] Simar Kareer, Naoki Yokoyama, Dhruv Batra, Sehoon Ha, and Joanne Truong. ViNL: Visual Navigation and Locomotion Over Obstacles. In *International Conference on Robotics and Automation (ICRA)*, 2023.
- [38] Donghyun Kim, Daniel Carballo, Jared Di Carlo, Benjamin Katz, Gerardo Bleedt, Bryan Lim, and Sangbae Kim. Vision aided dynamic exploration of unstructured terrain with a small-scale quadruped robot. In *2020 IEEE International Conference on Robotics and Automation (ICRA)*, pages 2464–2470. IEEE, 2020.
- [39] Ashish Kumar, Zipeng Fu, Deepak Pathak, and Jitendra Malik. Rma: Rapid motor adaptation for legged robots. In *Robotics: Science and Systems*, 2021.
- [40] Frédéric Large, Dizan Vasquez, Thierry Fraichard, and Christian Laugier. Avoiding cars and pedestrians using velocity obstacles and motion prediction. In *IEEE Intelligent Vehicles Symposium, 2004*, pages 375–379, 2004.
- [41] Joonho Lee, Jemin Hwangbo, Lorenz Wellhausen, Vladlen Koltun, and Marco Hutter. Learning quadrupedal locomotion over challenging terrain. *Science robotics*, 5(47):eabc5986, 2020.
- [42] Zhongyu Li, Xue Bin Peng, Pieter Abbeel, Sergey Levine, Glen Berseth, and Koushil Sreenath. Robust and versatile bipedal jumping control through reinforcement learning. In *Robotics: Science and Systems*, 2023.
- [43] Qingkai Liang, Fanyu Que, and Eytan Modiano. Accelerated primal-dual policy optimization for safe reinforcement learning. *arXiv preprint arXiv:1802.06480*, 2018.
- [44] Qiayuan Liao, Zhongyu Li, Akshay Thirugnanam, Jun Zeng, and Koushil Sreenath. Walking in narrow spaces: Safety-critical locomotion control for quadrupedal robots with duality-based optimization. In *2023 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2723–2730. IEEE, 2023.
- [45] Minghuan Liu, Menghui Zhu, and Weinan Zhang. Goal-conditioned reinforcement learning: Problems and solutions. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence*, 2022.
- [46] Yuntao Ma, Farbod Farshidian, and Marco Hutter. Learning arm-assisted fall damage reduction and recovery for legged mobile manipulators. In *2023 IEEE International Conference on Robotics and Automation (ICRA)*, pages 12149–12155. IEEE, 2023.
- [47] Viktor Makoviychuk, Lukasz Wawrzyniak, Yunrong Guo, Michelle Lu, Kier Storey, Miles Macklin, David Hoeller, Nikita Rudin, Arthur Allshire, Ankur Handa, et al. Isaac gym: High performance gpu based physics simulation for robot learning. In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*, 2021.
- [48] Gabriel B Margolis, Ge Yang, Kartik Paigwar, Tao Chen, and Pulkit Agrawal. Rapid locomotion via reinforcement learning. *arXiv preprint arXiv:2205.02824*, 2022.
- [49] Matias Mattamala, Nived Chebrolu, and Maurice Fallon. An efficient locally reactive controller for safe navigation in visual teach and repeat missions. *IEEE Robotics and Automation Letters*, 7(2):2353–2360, 2022.
- [50] Takahiro Miki, Joonho Lee, Jemin Hwangbo, Lorenz Wellhausen, Vladlen Koltun, and Marco Hutter. Learning robust perceptive locomotion for quadrupedal robots in the wild. *Science Robotics*, 7(62):eabk2822, 2022.
- [51] Yashwanth Kumar Nakka, Anqi Liu, Guanya Shi, Anima Anandkumar, Yisong Yue, and Soon-Jo Chung. Chance-constrained trajectory optimization for safe exploration and learning of nonlinear systems. *IEEE Robotics and Automation Letters*, 6(2):389–396, 2020.
- [52] NVIDIA Corporation. NVIDIA Warp repository. <https://github.com/NVIDIA/warp>.
- [53] Alex Ray, Joshua Achiam, and Dario Amodei. Benchmarking safe exploration in deep reinforcement learning. <https://openai.com/research/benchmarking-safe-exploration-in-deep-reinforcement-learning>, 2019.
- [54] Nikita Rudin, David Hoeller, Marko Bjelonic, and Marco Hutter. Advanced skills by learning locomotion and local navigation end-to-end. In *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2497–2503. IEEE, 2022.
- [55] Nikita Rudin, David Hoeller, Philipp Reist, and Marco Hutter. Learning to walk in minutes using massively parallel deep reinforcement learning. In *Conference on Robot Learning*, pages 91–100. PMLR, 2022.
- [56] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- [57] Mingyo Seo, Ryan Gupta, Yifeng Zhu, Alexy Skoutnev, Luis Sentis, and Yuke Zhu. Learning to walk by steering: Perceptive quadrupedal locomotion in dynamic environments. In *IEEE International Conference on Robotics and Automation (ICRA)*, 2023.
- [58] Young-Ha Shin, Tae-Gyu Song, Gwanghyeon Ji, and Hae-Won Park. Actuator-constrained reinforcement learning for high-speed quadrupedal locomotion. *arXiv preprint arXiv:2312.17507*, 2023.
- [59] Wenwen Si, Tianhao Wei, and Changliu Liu. Agen: Adaptable generative prediction networks for autonomous driving. In *2019 IEEE Intelligent Vehicles Symposium (IV)*, pages 281–286. IEEE, 2019.
- [60] Oswin So, Zachary Serlin, Makai Mann, Jake Gonzales, Kwesi Rutledge, Nicholas Roy, and Chuchu Fan. How to train your neural control barrier function: Learning

- safety filters for complex input-constrained systems. In *2024 IEEE International Conference on Robotics and Automation (ICRA) (Under Review)*, 2024.
- [61] Yunlong Song, Angel Romero, Matthias Müller, Vladlen Koltun, and Davide Scaramuzza. Reaching the limit in autonomous racing: Optimal control versus reinforcement learning. *Science Robotics*, 8(82):eadg1462, 2023.
 - [62] Stereolabs Inc. Fill mode - depth settings. <https://www.stereolabs.com/docs/depth-sensing/depth-settings#fill-mode>, 2024.
 - [63] ShangJie Sun, ShuHai Jiang, SongHe Cui, Yue Kang, YuTang Chen, et al. Path planning of forest fire-fighting robots based on deep learning. *Forest Engineering*, 36(4):51–57, 2020.
 - [64] Sangli Teng, Yukai Gong, Jessy W. Grizzle, and Maani Ghaffari. Toward safety-aware informative motion planning for legged robots. *CoRR*, abs/2103.14252, 2021.
 - [65] Chen Tessler, Daniel J Mankowitz, and Shie Mannor. Reward constrained policy optimization. *arXiv preprint arXiv:1805.11074*, 2018.
 - [66] Brijen Thananjeyan, Ashwin Balakrishna, Suraj Nair, Michael Luo, Krishnan Srinivasan, Minho Hwang, Joseph E Gonzalez, Julian Ibarz, Chelsea Finn, and Ken Goldberg. Recovery rl: Safe reinforcement learning with learned recovery zones. *IEEE Robotics and Automation Letters*, 6(3):4915–4922, 2021.
 - [67] Josh Tobin, Rachel Fong, Alex Ray, Jonas Schneider, Wojciech Zaremba, and Pieter Abbeel. Domain randomization for transferring deep neural networks from simulation to the real world. In *2017 IEEE/RSJ international conference on intelligent robots and systems (IROS)*, pages 23–30. IEEE, 2017.
 - [68] Marco Tranzatto, Takahiro Miki, Mihir Dharmadhikari, Lukas Bernreiter, Mihir Kulkarni, Frank Mascarich, Olov Andersson, Shehryar Khattak, Marco Hutter, Roland Siegwart, et al. Cerberus in the darpa subterranean challenge. *Science Robotics*, 7(66):eabp9742, 2022.
 - [69] Joanne Truong, Max Rudolph, Naoki Harrison Yokoyama, Sonia Chernova, Dhruv Batra, and Akshara Rai. Rethinking sim2real: Lower fidelity simulation leads to higher sim2real transfer in navigation. In *Conference on Robot Learning*, pages 859–870. PMLR, 2023.
 - [70] Joanne Truong, April Zitkovich, Sonia Chernova, Dhruv Batra, Tingnan Zhang, Jie Tan, and Wenhao Yu. Indoorsim-to-outdoorreal: Learning to navigate outdoors without any outdoor experience. In *arXiv preprint arXiv:2305.01098*, 2023.
 - [71] USA Today. Robots taking on tasks from mundane to dangerous: Police robot dog shot by suspect. <https://www.usatoday.com/story/tech/2024/03/28/what-to-know-about-spot-the-robotic-dog/73118781007/>, 2024.
 - [72] Yikai Wang, Mengdi Xu, Guanya Shi, and Ding Zhao. Guardians as you fall: Active mode transition for safe falling. *arXiv preprint arXiv:2310.04828*, 2023.
 - [73] Lorenz Wellhausen and Marco Hutter. Artplanner: Robust legged robot navigation in the field. *Field Robotics*, 3(1):413 – 434, 2023-03. ISSN 2771-3989. doi: 10.3929/ethz-b-000614683.
 - [74] Erik Wijmans, Manolis Savva, Irfan Essa, Stefan Lee, Ari S. Morcos, and Dhruv Batra. Emergence of maps in the memories of blind navigation agents. In *International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=ITt4KjHSsYl>.
 - [75] Wenli Xiao, Tairan He, John Dolan, and Guanya Shi. Safe deep policy adaptation. *arXiv preprint arXiv:2310.08602*, 2023.
 - [76] Chuanyu Yang, Kai Yuan, Qiuguo Zhu, Wanming Yu, and Zhibin Li. Multi-expert learning of adaptive legged locomotion. *Science Robotics*, 5(49):eabb2174, 2020.
 - [77] Ruihan Yang, Minghao Zhang, Nicklas Hansen, Huazhe Xu, and Xiaolong Wang. Learning vision-guided quadrupedal locomotion end-to-end with cross-modal transformers. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=nhnJ3oo6AB>.
 - [78] Tsung-Yen Yang, Tingnan Zhang, Linda Luu, Sehoon Ha, Jie Tan, and Wenhao Yu. Safe reinforcement learning for legged locomotion. In *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2454–2461. IEEE, 2022.
 - [79] Yuxiang Yang, Guanya Shi, Xiangyun Meng, Wenhao Yu, Tingnan Zhang, Jie Tan, and Byron Boots. Cajun: Continuous adaptive jumping using a learned centroidal controller. *arXiv preprint arXiv:2306.09557*, 2023.
 - [80] Naoki Yokoyama, Alex Clegg, Joanne Truong, Eric Undersander, Tsung-Yen Yang, Sergio Arnaud, Sehoon Ha, Dhruv Batra, and Akshara Rai. Asc: Adaptive skill coordination for robotic mobile manipulation. *IEEE Robotics and Automation Letters*, 9(1):779–786, 2023.
 - [81] Chong Zhang, Wanming Yu, and Zhibin Li. Accessibility-based clustering for efficient learning of locomotion skills. In *2022 International Conference on Robotics and Automation (ICRA)*, pages 1600–1606. IEEE, 2022.
 - [82] Chong Zhang, Nikita Rudin, David Hoeller, and Marco Hutter. Learning agile locomotion on risky terrains. *arXiv preprint arXiv:2311.10484*, 2023.
 - [83] Chong Zhang, Jin Jin, Jonas Frey, Nikita Rudin, Matias Eduardo Mattamala Aravena, Cesar Cadena, and Marco Hutter. Resilient legged local navigation: Learning to traverse with compromised perception end-to-end. In *41st IEEE Conference on Robotics and Automation (ICRA 2024)*, 2024.
 - [84] Weiye Zhao, Tairan He, and Changliu Liu. Model-free safe control for zero-violation reinforcement learning. In *5th Annual Conference on Robot Learning*, 2021.
 - [85] Weiye Zhao, Tairan He, Rui Chen, Tianhao Wei, and Changliu Liu. State-wise safe reinforcement learning: A survey. In Edith Elkind, editor, *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence, IJCAI-23*, pages 6814–6822. International Joint

Conferences on Artificial Intelligence Organization, 8
2023. URL <https://doi.org/10.24963/ijcai.2023/763>.

- [86] Ziwen Zhuang, Zipeng Fu, Jianren Wang, Christopher Atkeson, Soeren Schwertfeger, Chelsea Finn, and Hang Zhao. Robot parkour learning. *arXiv preprint arXiv:2309.05665*, 2023.

APPENDIX

A. Training Costs

On an NVIDIA RTX 4090, each iteration of the agile policy learning takes 1.5 sec. The π^{Agile} policy needs ~ 800 iterations (~ 20 min) to converge, and the LAG policy needs ~ 2500 iterations (~ 1 h) to converge. Yet, to fully unleash their performance and make a fair comparison, all of the policies are tested with the checkpoints of 10000 iterations (~ 4 h). The recovery policy is far more efficient and converges within 500 iterations, becoming near-optimal within 10 min. The GPU memory usage during training is below 5GB.

We train the RA value network in parallel with the agile policy rollout to save time, and the procedure takes ~ 40 min.

The ray-prediction network can use collected data from different trained policies as the prediction in ideal cases is not conditioned on the policy. The network training time can vary for different settings. As a reference, we collected 250k labelled images for training.

B. Obtaining Ray Distances in Simulation

Ideally, we should use ray tracing tools such as Embree [33] and Warp [52] to obtain accurate ray distances. However, these tools are time-consuming and GPU-intensive for our separated parallel environments with multiple objects. To enable efficient training, we implement an analytical ray tracing calculation, which, while efficient, is limited to simple geometries such as cylinders. Therefore, we are unable to use diverse objects during training, and have to choose the objects with meshes that can be approximated during ray-prediction data collection.

Despite this limitation, our design of the low-dimensional exteroception representation, together with domain randomization and data augmentation, enables the system to generalize to a variety of real-world objects, as demonstrated in the Figure 1.

C. RL Hyperparameters

TABLE X
RL HYPERPARAMETERS

Hyperparameter	Value
Agile Policy	
#. steps per iteration	48
Entropy coefficient	0.003
Actor NN	MLP with hidden units [512, 256, 128]
Critic NN	MLP with hidden units [512, 256, 128]
Others	Same as [55]
Agile Policy - LAG	
Cost-critic NN	MLP with hidden units [512, 256, 128]
Cost-critic loss coefficient	1.0
Cost gamma	0.99
Cost lambda	0.97
Cost limit	0.0
Multiplier learning rate	0.001
Recovery Policy	
#. steps per iteration	24
Entropy coefficient	0.003
Actor NN	MLP with hidden units [512, 256, 128]
Critic NN	MLP with hidden units [512, 256, 128]
Others	Same as [55]

D. Explanations for Regularization Terms

We explain why we add so many regularization terms in Table XI, which follows [55, 54, 48, 82, 29] and has some adaptations based on our tuning experiences.

TABLE XI
EXPLAINING REGULARIZATION REWARDS

Term	Reason
v_z^2	To reduce vertical oscillation
$\omega_x^2 + \omega_y^2$	To reduce rotational oscillation
$(g_x^2 + g_y^2)$	To reduce tilting
$\ \tau\ _2^2$	To reduce mechanical stress and power consumption of motions
$\sum_{i=1}^{12} \text{ReLU}(\tau_i - 0.85 \cdot \tau_{i,\text{lim}})$	To avoid large torques that can lead to sim-to-real gaps
$\ \dot{q}\ _2^2$	To reduce aggressive motions that may damage the hardware
$\sum_{i=1}^{12} \text{ReLU}(\dot{q}_i - 0.9 \cdot \dot{q}_{i,\text{lim}})$	To avoid large joint velocities that can lead to inaccurate simulation
$\sum_{i=1}^{12} \text{ReLU}(q_i - 0.95 \cdot q_{i,\text{lim}})$	To avoid joint positions near the limit that can lead to inaccurate simulation
$\ \ddot{q}\ _2^2$	To reduce jerky motions
$\ \dot{a}\ _2^2$	To encourage smooth actions
$\mathbb{1}(\text{fly})$	To constrain #. uncontrollable DoFs and improve maneuverability