

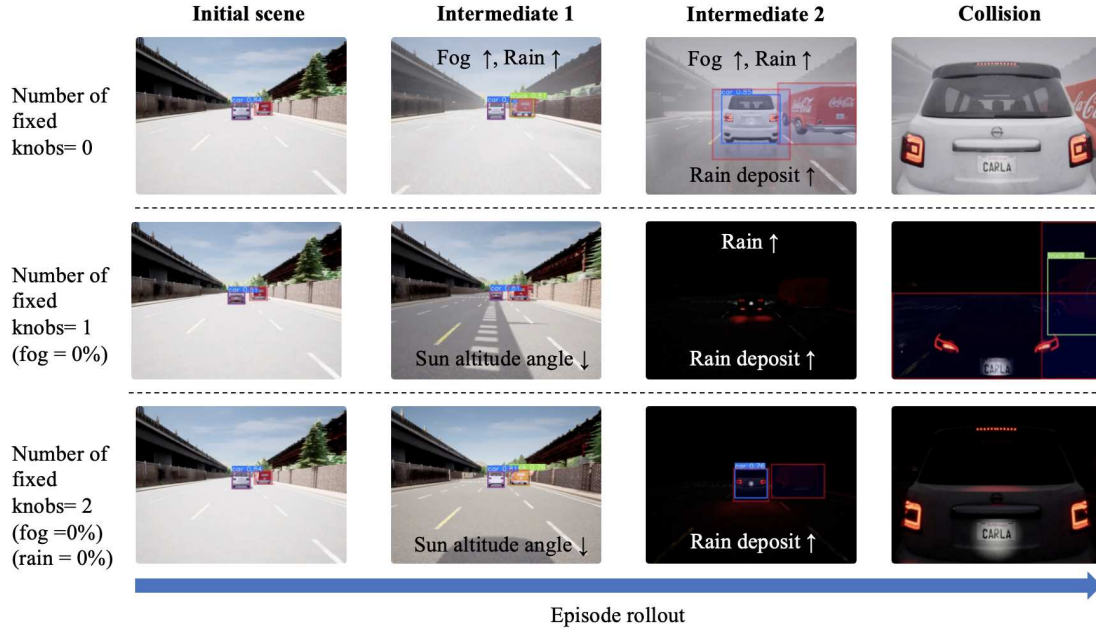
GENESIS-RL: GEnerating Natural Edge-cases with Systematic Integration of Safety considerations and Reinforcement Learning

Hsin-Jung Yang*, Joe Beck[†], Md Zahid Hasan*, Ekin Beyazit*
Subhadeep Chakraborty[†], Tichakorn Wongpiromsarn*, Soumik Sarkar*

*Iowa State University, Ames, IA, USA

[†]University of Tennessee, Knoxville, TN, USA

Email: {hgy, zahid, ekin, nok, soumiks}@iastate.edu, {jbeck9, schakrab}@utk.edu



GENESIS-RL: A reinforcement learning framework to progressively manipulate the environment (weather conditions in this example) for an autonomous system (autonomous vehicle in this case) to systematically synthesize natural edge cases that may lead to system-level safety issues (collision in this case). Project website: <https://hgy77.github.io/GENESIS-RL/>

Abstract—In the rapidly evolving field of autonomous vehicles, the safety and reliability of the system components are fundamental requirements. These components are often vulnerable to complex and unforeseen environments, making natural edge-case generation essential for enhancing system resilience. This paper presents GENESIS-RL, a novel framework that leverages system-level safety considerations and reinforcement learning techniques to systematically generate naturalistic edge cases. By simulating challenging conditions that mimic the real-world situations, our framework aims to rigorously test entire system’s safety and reliability. Our experimental validation, conducted on high-fidelity simulator underscores the overall effectiveness of this framework.

I. INTRODUCTION

Scenario-based testing is one of the key approaches for the validation of autonomous systems, especially those that

incorporate learning-enabled components that are known to be susceptible to rare, unexpected (potentially out-of-distribution) scenarios [1]. This testing approach is vital not only for ensuring the safety and reliability of these systems but also for enabling them to identify and rectify potential failures in diverse, unforeseen situations. In this context, identifying and preparing for challenging or edge-case scenarios becomes critical. Synthesizing realistic edge-case samples and incorporating these into the training process [2], [3] can significantly enhance the resilience of learning-enabled modules against adversarial conditions. By exposing the learning modules to these pessimistic samples, systems gain the opportunity to learn from challenging data and better generalize across a spectrum of real-world conditions. However, given the vast amount of possible scenarios, manual creation of every scenario is infeasible, making automated edge-case generation crucial for scalability and effectiveness [4].

This work was partly supported by the National Science Foundation, USA under grants CNS-1845969, CNS-2141153, CNS-1954556.

In this regard, traditional adversarial attacks on machine learning models explore the vulnerability of the models by injecting imperceptible noise into the inputs [5], [6]. These input perturbation methods, while effective in degrading the model performance, typically generates unnatural and unrealistic samples, diverging from genuine real-world scenarios. Furthermore, these approaches usually target specific components of an autonomous system rather than assessing the system as a whole. This narrow focus can overlook the holistic behavior of the system, where, for instance, a failure in the perception module might be compensated by the system’s control mechanisms, thus not leading to a failure at the system level. On the other hand, if the control system is not able to compensate, a relatively small error in the perception module may lead to a catastrophic system-level failure. This highlights the limitation of focusing solely on component-level vulnerabilities without considering the integrated operation of the entire system.

Generative models have been used to synthesize edge cases that are more realistic [7]. However, they are known to produce samples with artifacts that compromise their realism. These models, including generative adversarial networks (GANs) [8], diffusion models [9], and more recently, text-to-image generation models such as DALL-E [10], CogView [11], can suffer from issues such as unnatural distributions [12], distinct artifacts and unstable training [13], [14], and slow inference rates, limiting their effectiveness in producing realistic and natural scenarios [15], [16].

In this paper, we aim to alleviate these challenges, by performing edge-case generation with system-level safety objectives while maintaining the naturalness of the generated scenarios. We employ the rulebook formalism [17] to precisely specify system-level safety objectives and leverage the capabilities of Reinforcement Learning (RL) to guide the generation of scenarios that not only challenge the system across all its components but also resemble real-world conditions closely. By focusing on the end-to-end vulnerability of autonomous systems, our approach aims to generate scenarios where the system fails to adhere to rulebook safety rules, thereby identifying potential systemic failures. Also, our proposed framework ensures that the generated scenarios are not only challenging but also devoid of unrealistic artifacts (via use of high-fidelity simulators), offering a more effective and comprehensive approach to testing and validating the safety and reliability of autonomous systems.

In summary, the key contributions of this paper are as follows:

- We propose a synthetic edge case data generation framework for system-level safety concerns in learning-enabled autonomous systems.
- We propose an RL-based intelligent sampling technique that can identify parametric settings of high-fidelity simulators to generate natural edge cases that may lead to violation of safety rules by a learning-enabled autonomous system.
- We provide extensive experimental validation of our framework using the CARLA simulator [18]. We also demonstrate that a pre-trained RL policy can generate edge-cases for new scenarios with minimal to no training steps, thus accelerating the process of assessment and verification of learning-enabled autonomous systems.

II. RELATED WORKS

Recent research has explored diverse approaches to generating edge cases. Efforts using cost functions to pinpoint high-risk scenarios have shown potential yet often neglect critical factors like unpredictable trajectories [19]–[21]. Perception-based techniques, such as constant norm-based perturbation, target the system’s perception capabilities but may not address the system’s overall performance comprehensively [2], [7], [22], [23]. While innovative, methods that extract and recreate accidents from videos face challenges in accurately replicating real-world complexity [24].

Additionally, some edge-case generation software toolkits, like VerifAI [25], are capable of analysis, falsification, and data augmentation for systems utilizing ML architectures. These toolkits leverage an “abstract feature space” of higher-level information compared to the low-level “concrete feature space” of image pixels to search for rule violation scenarios in a given environment. Domain randomization effectively bridges the sim-to-real gap [26], [27], but it can lead to training an overly conservative policy depending on the range of randomization. System identification [28] offers a feasible solution by estimating environmental parameters through limited interaction with real-world scenarios.

Lastly, Bayesian optimization-based methods often generate challenging scenarios with limited diversity and insufficient complexity [2], [19]. These methods typically produce short scenario segments with limited environment interactions since they require a predefined parameter range [29]. Consequently, they limit the assessment of system performance and fail to capture realistic edge cases with diverse interactions.

III. BACKGROUND

In this work, a *system* refers to the entity that is being evaluated for its ability to navigate and perform tasks within variable conditions. It could be an autonomous vehicle or any computational model. The *world* denotes the simulated surroundings in which the system operates, a construct designed to emulate real-world dynamics where every aspect can have an effect on the system’s behavior. Lastly, an *actor* is an entity other than the system that also lives in the world.

As an example, in an autonomous driving context, the system could be the ego vehicle, and the world is where the ego vehicle is situated. Other entities, such as other vehicles and pedestrians on the street, are actors. Together with crucial factors such as weather and road conditions (including road markings and traffic signs) that are not part of the system but could affect the system’s behavior, they are all parts of the world.

A. Deep Reinforcement Learning

Deep Reinforcement Learning (DRL) [30] is an extension of RL that harnesses the representational power of deep neural networks. At its core, DRL adheres to the Markov Decision Process (MDP) framework, mathematically formulated as a 4-tuple (S, A, P, R) , where:

- S represents the state space, comprising all conceivable states s_t at a given time t .
- A denotes the action space, encompassing all actions a_t available to the agent at time t .
- P the transition probability function, indicating the probability of transitioning from one state s_t to another state s_{t+1} given an action a_t .
- $R : S \times A \times S' \rightarrow \mathbb{R}$ is the reward function, which assigns a numerical reward for each transition between states under specific actions.

In a typical DRL setup, *the DRL agent* is the entity that we hope to train, whereas *the environment* is the setting or domain wherein the DRL agent operates and makes decisions, which encompasses all aspects mentioned above in the MDP framework, including the state space S , action space A , the transition probability function P and rewards R .

Under this MDP framework, the agent’s decision-making strategy at any time t is governed by a policy π , which maps the current state s_t to an action a_t . In DRL, this policy is represented with neural networks, denoted as π_θ , where θ represents the neural network’s trainable parameters. This configuration enables the agent to dynamically refine its strategy by updating θ , thus improving its performance and adaptability in navigating the environment.

The objective of DRL is to discover an optimal policy π_θ^* that guides the agent to maximize the expected return along a trajectory τ , which is a sequence of states and actions $(s_0, a_0, s_1, a_1, \dots, s_T, a_T)$. The expected return is calculated as $J(\theta) = \mathbb{E}_{\tau \sim \pi_\theta} [\sum_{t=0}^T \gamma^t R(s_t, a_t, s_{t+1})]$, where γ is the discount factor and T the length of the trajectory. The concept of episodes emerges naturally from this setup. An episode describes a complete trajectory from an initial state to a terminal state [31].

B. Rulebook

We will use the rulebook formalism [17] to precisely describe the correct behavior of the system. A rulebook consists of a set Λ of rules; each is evaluated over realizations. A realization is defined as a sequence of states of the system and all the other actors and features in the world. Given a set Σ of realizations, a rule is defined as a function $\lambda : \Sigma \rightarrow \mathbb{R}_{\geq 0}$ that measures the degree of violation of its argument. If $\lambda(x) < \lambda(y)$, then the realization y violates the rule λ to a greater extent than does x . In particular, $\lambda(x) = 0$ indicates that a realization x is fully compliant with the rule. Note that the definition of the violation metric might be analytical, “from first principles”, or be the result of a learning process.

In this work, we utilize the rulebook to calculate the rewards. A higher violation score leads to increased rewards,

encouraging the agent to explore scenarios that challenge the system’s safety protocols and resilience, thereby generating critical edge cases. For a detailed description of the rules used for reward calculation, please refer to the *Reward calculator* subsection in the *Experiments* section.

IV. METHODOLOGY

At a high level, GENESIS-RL utilizes DRL to dynamically explore and manipulate the conditions of a simulated world, aimed at generating challenging yet naturalistic edge-cases for a system. To achieve this, we parameterized the world with *parametric knobs*—adjustable settings that control various aspects of the simulation, which in the case of autonomous driving, could include dynamic weather patterns, object placements, traffic flow, and so on. By adjusting these knobs, the DRL agent is provided with the capability to systematically probe and alter the simulated world, effectively simulating different edge cases that the system under test might encounter.

Remark. *Our objective is to craft and manipulate the world (via simulation) to induce challenging scenes. By doing so, we seek to generate edge cases that test the limits of the system’s current capabilities, aiming to reveal potential failure cases. In contrast to typical DRL works, we do not focus on improving the system’s capabilities in this work.*

A. DRL Problem Formulation

Following the MDP framework, we define the state space, action space, and reward of our problem as follows:

1) *State space*: The state space encompasses all conceivable states s_t , including permutations of parametric knobs, the system’s behaviors, other actors, and features of the world. This state representation captures the dynamics of the world and the DRL agent’s action inputs, and is conveyed through information obtained by the system.

2) *Action space*: The action space is the set of all possible actions a_t available to the agent, corresponding to the adjustments the agent can make to the parametric knobs within the simulation. To ensure that the changes introduced by the DRL agent lead to scenes that are natural and realistic, we imposed constraints on the extent of modifications possible at each step. Specifically, we limit the maximum percentage change that can be applied to any parametric knob by the DRL agent in a single action. This measure prevents extreme, unrealistic variations in conditions, thereby maintaining the realistic nature of the simulated scenes while still challenging the system under test.

3) *Reward*: The reward mechanism is designed to motivate the DRL agent to discover edge cases. It comprises two components: the learning module loss r_m and the violation score r_v derived from the rulebook. The learning module loss is the loss experienced by the learning-enabled module within the system, which acts as a direct reward to the agent, where a lower loss indicates better performance of the module at performing its designated task. The violation score is an indirect reward provided to the agent due to the imperfection of the learning-enabled modules. For example, in autonomous driving, the rulebook evaluates the ego vehicle’s trajectory

against a set of predefined rules, penalizing actions that lead to unsafe scenarios. The total reward r_t at time step t is calculated as a combination of these two elements.

B. GENESIS-RL Framework

To implement our DRL formulation, we designed a framework consisting of the following components: the DRL agent, the initial scene generator, the simulator, the system, and the reward calculator. The latter four together form an environment for the DRL agent, facilitating continuous learning of the DRL agent through dynamic interaction.

1) *DRL agent*: The DRL agent is the decision-making core. At each time step t , it obtains the current state s_t of the environment and executes an action a_t . The environment then responds to this action by evolving to a new state based on the updated parametric knobs of the simulated world and issues a scalar reward r_t to the agent as feedback.

2) *Initial scene generator*: The initial scene generator is responsible for creating a distribution of the initial scenes (a configuration of physical objects, the system and actors) and sampling from them in the simulated world. It dictates the initial conditions the system will encounter, therefore determining the initial scene observed by the DRL agent.

3) *Simulator*: The simulator provides a realistic and interactive backdrop where the DRL agent's actions and the system's outputs are executed and new frames are updated, reflecting the changes in real-time.

4) *System*: As defined in the background section, The system is the entity being evaluated within variable conditions.

5) *Reward calculator*: As defined in the previous section, the reward calculator calculates the reward r_t for time step t .

C. Training the DRL Agent

Putting things together, a single step of the training looks like as follows (See Fig. 1): at each time step t , the DRL agent receives a state s_t from the simulator and executes an action a_t on the simulator. The simulator reflects the changes based on the updated parametric knobs and the changes are subsequently captured by the system through its sensors. The system then generates control signals based on its inputs, which leads to a system trajectory update. The updated trajectory is then evaluated by the rulebook for violation score calculation and is sent back to the DRL agent combined with the learning module loss.

The DRL agent is trained through interactions with the environment, where it observes the states, applies actions, and receives rewards. The training process involves iterative episodes of simulation, during which the agent refines its policy π_θ to maximize the cumulative reward, effectively learning to identify and create challenging scenarios for the system.

V. EXPERIMENTS

In this paper, we explore the weather conditions that can lead to natural edge cases for autonomous driving. Hence, we grant the DRL agent exclusive control over the weather

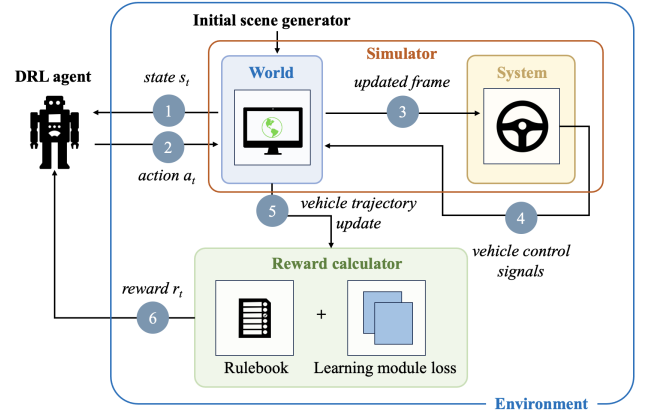


Fig. 1. Architectural overview of the proposed framework. At each step t , the DRL agent observes a state s_t (1) and executes an action a_t (2). The simulator then updates the simulated world accordingly and creates an updated frame. The updated frame (3) is then processed by the system to generate vehicle control signals (4). The control signals are then applied to the simulated world to update the vehicle trajectory (5). The reward calculator evaluates the performance by comparing the ego vehicle's trajectory against the rulebook and also computes the learning module loss, issuing a scalar reward r_t (6) that guides the DRL agent's learning process.

conditions in the simulated world. The system we evaluate is the ego vehicle, tasked to navigate through the simulated world based on sensor feedback. In subsequent sections, we detail the operationalization of the GENESIS-RL framework's components, starting with the DRL agent and encompassing the environment, then the training and testing setups and evaluation metrics.

Remark. *Despite only having weather parameters as available actions in this work, we can also include other factors such as the behavior of other actors into our framework by parameterizing their behavior in the simulation (such as the aggressiveness of the other actors using CARLA's traffic manager, or the position and orientation of the actors using Scenic [32]).*

A. DRL Agent

We have chosen the Proximal Policy Optimization (PPO) algorithm [33] as the DRL agent for its stabilized training capabilities and proficiency in handling continuous state and action spaces. We adopted the implementation of Stable-Baselines3 [34]. The state and action spaces are as follows:

- **State space:** A 640×480 three channel numpy array from the RGB camera attached to the ego vehicle.
- **Action space:** A 5-tuple that represents the delta in value for each parametric knob. The parametric knobs are: the fog density, the precipitation density, the precipitation deposit level, the sun altitude angle, and the sun azimuth angle. These actions are bounded between $[-1, 1]$, and are scaled by a set of preset scalar factors to satisfy the perturbation limit, which is 5%. For example, if the fog density ranges from 0 to 100, then the corresponding

parametric knob would be scaled by a factor of 5 such that the change in fog density level between each step in an episode will not exceed 5%.

As mentioned earlier, the DRL agent is developing a policy π_θ parameterized by a neural network θ . Here, the architecture of θ is a CNN, which takes the RGB three channel arrays (640, 480, 3) as inputs, and outputs a 5-tuple action.

B. Environment

The initial scene generator, the simulator, the system, and the reward calculator collaboratively create the environment of the DRL agent. These parts together generate the state and reward required to train the DRL agent and react to the actions of the DRL agent dynamically.

1) *Simulator*: We employ CARLA [18] to simulate intricate, dynamic urban settings with high visual fidelity.

2) *Initial scene generator*: Scenic [32] was utilized for creating the initial scenes as a strategic approach to ensure precision and versatility in our experimental setup. It allows us to craft realistic, detailed and specific initial scenes, thereby enhancing the relevance and challenge of each test instance presented to the system. Specifically, this generator is tasked with assigning the positions of vehicles within the simulator: For each initial scene, we introduce 30 vehicles, including the ego vehicle, a lead vehicle, a neighboring vehicle on the front right, and the remaining 27 vehicles randomly positioned within a specific radius of the ego vehicle. Vehicle colors, models, and the distance between the ego, lead, and neighbor vehicles could be either varied randomly or deterministic, depending on the use case. All vehicles except for the ego vehicle are set to CARLA's autopilot in order to provide a dynamic environment.

We leveraged CARLA's [18] record feature to efficiently catalog the configurations of the initial scenes into json files. These saved scenes are then randomly sampled and loaded back into the simulation between training or testing episodes. A total of 1199 initial scenes for *Town05* and 1202 for *Town10* were generated and cataloged using this generator.

3) *System*: The system in our study is an ego vehicle that navigates itself within the simulated world with RGB and depth cameras onboard. It is equipped with a perception-based controller, which consists of two parts: a perception model for object detection, utilizing a pre-trained YOLOv5s model [35] to process RGB images from the onboard RGB camera, and a modified CARLA [18] behavior controller that translates detection results into vehicle control signals. In our implementation, the input to the modified behavior controller includes the detection output from the perception model along with depth information from the onboard depth camera. By overlapping the 2D bounding boxes with the depth image, we were able to obtain the depth of the center point of the bounding boxes, which will then be used to obtain the 3D bounding boxes of the detected objects. Leveraging these 3D bounding boxes, combined with real-time data on the vehicle's position, orientation, velocity, and the surrounding map, the

controller adeptly synthesizes these inputs to generate vehicle control signals, i.e., throttle, brake, and steering signals.

4) *Reward calculator*: The reward r_t is calculated by the weighted sum of the learning module loss r_m and the violation score r_v . The learning module loss is defined by the Intersection Over Union (IOU) metric, assessing the system's object detection precision in real-time. Simultaneously, the violation score is generated by the rulebook evaluating a trajectory $\tau = s_1, s_2, s_3, \dots, s_T$, consisting of the following two violation scores based on two critical safety rules:

- Vehicle collision rule, $\lambda_c(\tau)$, violated if the ego vehicle collides with other vehicles:

$$\lambda_c(\tau) = \sum_{t=1}^T \alpha_c(s_t) \quad (1)$$

$$\alpha_c(s_t) = v_c \text{ if collision happens, } 0 \text{ otherwise} \quad (2)$$

- Vehicle proximity rule, $\lambda_p(\tau)$, violated when the distance d between the ego vehicle and the vehicle in front is less than a predetermined distance (which is set to 5 meters):

$$\lambda_p(\tau) = \sum_{t=1}^T \alpha_p(s_t) \quad (3)$$

$$\alpha_p(s_t) = v_p \text{ if } d < 5, 0 \text{ otherwise} \quad (4)$$

where v_c and v_p are the speeds of the ego vehicle in meters per second when a violation occurs. $\alpha_c(s_t)$ and $\alpha_p(s_t)$ are then weighted by their respective weights w_c and w_p then summed and transformed using the natural logarithm to compute r_v . Finally, the reward r_t is calculated as follows:

$$r_t = r_m + r_v = e^{-iou} + \ln(1 + w_c \alpha_c + w_p \alpha_p) \quad (5)$$

where iou is the IOU metric, $(w_c, w_p) = (500, 100)$. The rationale for incorporating speed values when a violation occurs is to ensure that violations occurring at higher speeds are assigned higher violation scores.

C. Training and Testing

Our experiments are divided into two phases: training and testing, each aimed at evaluating the effectiveness of our framework.

1) *Training*: The training experiments were conducted in CARLA's [18] *Town10*, an urban map setting with multiple four-way intersections. To ensure the DRL agent experiences a broad range of initial scene configurations, we start each episode by randomly selecting from 1202 pre-cataloged scenes, each featuring vehicles with randomized colors and makes. This setup is further enhanced by the stochastic behaviors of non-ego vehicles, governed by CARLA's [18] traffic manager, introducing both realism and complexity. The variable vehicle arrangement, appearances, and the other vehicle agents' stochastic nature significantly enriched the training landscape, equipping the DRL agent to adeptly uncover the system's vulnerabilities. The DRL agent is trained with 40960 steps, with an episode length of 512. The policy of the DRL agent updates every four trajectories.

2) *Testing*: The testing experiments unfolded within CARLA’s [18] *Town05*, a map setting featuring pine-covered hills and a network of roads and a highway. In this phase, we aimed for diverse yet repeatable conditions by selecting 50 out of 1199 cataloged scenes from the initial scene generator. This is realized by using a fixed random seed, ensuring these scenarios are consistent across different experiment runs. The deterministic setting extends to other vehicles’ appearances and their behavior, i.e., the color/make of the vehicles and the path they follow throughout the testing episode are deterministic, chosen to maintain uniformity to allow each scenario’s outcomes to be directly comparable. This method ensures that despite the inherent variability of the test runs, the foundational conditions remain constant, facilitating an accurate assessment of performance across varying scenarios.

Records of violation scores and other information were kept, with averages calculated for analysis. These findings were benchmarked against two specific scenarios: our system navigating in clear weather and under randomly perturbed weather conditions by a non-strategic agent, providing a comprehensive evaluation of the GENESIS-RL framework. The length of each testing episode is identical to that of the training, which is 512 time steps in the simulation.

D. Evaluation Metrics

In our evaluation process, we employ two metrics: the violation score and the minimum following distance deficit δ_{mfd} . The former metric has been detailed earlier in this section and measures the system’s adherence to a set of predefined safety rules. For the latter, we introduce a metric by leveraging the concept of minimum following distance d_{min} , as defined in Responsibility-Sensitive Safety (RSS) [36]. We adopt a simplified version of the RSS criterion, which assumes that both the ego and lead vehicles have the same maximum deceleration rate and that the reaction time of the autonomous system is negligible. The simplified formula for calculating d_{min} is given by the following equation:

$$d_{min} = \max\{0, (\frac{v_e^2 - v_l^2}{2a})\} \quad (6)$$

where v_e , v_l the speed of the ego and lead vehicle, respectively; a the maximum deceleration of both vehicles, which is 5 m/s^2 .

The minimum following distance deficit δ_{mfd} is then calculated by the discrepancy between the actual distance d maintained by the ego vehicle and the calculated d_{min} , i.e., $\delta_{mfd} = \max(0, d_{min} - d)$, where d is the distance between the ego and lead vehicle. Specifically, it quantifies how much closer the ego vehicle comes to the lead vehicle than is deemed safe according to the simplified RSS criterion. In the following section, we will show the sum of this metric over all 50 runs. Showing the extent of the ego vehicle getting too close to the lead car.

E. Computational Time

The experiments were conducted on a computational setup with a 32-core CPU, 64GB of RAM, and an NVIDIA GTX

TITAN X graphics card. A complete training run required approximately 3-4 hours on this hardware, while a testing run took about 1-2 minutes per episode.

VI. RESULTS AND DISCUSSION

In this section, we present the outcomes of our experiments, including an analysis of the performance and efficacy of our framework. The results underscore the capability of our approach to generate meaningful edge cases, highlighting key findings and insights gained through testing scenarios. We validate our approach by answering the following questions:

- What impact do the generated edge cases have on the performance and decision-making processes of the system?
- How effectively does our framework generate edge cases that are both challenging and realistic for the system under test?

A. Impact of GENESIS-RL Generated Edge Cases on System Performance

Our analysis of GENESIS-RL’s impact on the system reveals a notable increase in both violation scores (See Fig. 2) as well as the sum of minimum following distance deficit, illustrating its ability to effectively challenge and exploit system vulnerabilities. Specifically, the scenario output from GENESIS-RL’s policy leads to significant variations in system performance, as demonstrated by the following observations:

1) *Proximity violation score increase*: Under weather scenarios controlled by GENESIS-RL, the system’s braking response was markedly delayed compared to its reaction under other testing scenarios and, in some instances, the braking behavior is altogether absent. This delay/absence is quantifiably demonstrated through the analysis of the ego vehicle’s telemetry data (See Fig. 3(c), the system operating under the GENESIS-RL policy maintains brakes much later, i.e., the blue curve is shifted more to the right compared to the two other scenarios), showcasing a contrast in the system’s ability to maintain safe following distances under varied environmental influences.

2) *Collision violation score increase*: Further looking into the types of collision violations induced by GENESIS-RL, we categorize them into three main failure modes:

- Non-detection collisions: Instances where the system fails to detect the leading vehicle at all, resulting in high-speed impacts (See Fig. 3(a)). This failure mode underscores critical perception system vulnerabilities under complex environmental conditions orchestrated by GENESIS-RL.
- Intermittent detection collisions: Occurrences where the system initially detects the lead vehicle but subsequently loses track of it, leading to collisions at reduced speeds (See Fig. 3(b)). These incidents underscore the deficiencies in the system’s ongoing tracking, and/or highlight the shortcomings in its ability to respond promptly within the scenarios induced by GENESIS-RL.
- Delayed detection collision: Occurrences where the system detects the lead vehicle too late to stop in time,

leading to collisions at reduced speeds (See Fig. 3(c)). These incidents highlight the deficiencies in the system’s detection mechanisms.

These findings not only exemplify GENESIS-RL’s capability in uncovering and leveraging system weaknesses but also emphasize the imperative need for bolstered system robustness against a wide spectrum of real-world conditions.

B. Effectiveness of GENESIS-RL in Generating Edge Cases

Our results show that with the GENESIS-RL framework, we were able to generate edge cases that pose significant challenges not only to automated systems but also to human perception and response capabilities. A prime example of such conditions includes scenarios combining foggy weather with heavy rainfall or nocturnal settings accentuated by heavy rain, where the reflections on wet surfaces severely disrupt the detection capabilities of autonomous systems. These conditions, inherently challenging due to their impact on visibility and sensor efficiency, highlight the scenario generation capabilities of GENESIS-RL, underscoring its potential for creating diverse testing environments that closely mimic real-world driving complexities.

During inference, we observed that the DRL agent relies heavily on manipulating certain parametric knobs (such as rain and fog density) to introduce challenges to the system. To investigate GENESIS-RL’s capability to generate a variety of edge cases, we conducted experiments with fog and rain density levels set to zero, individually and in combination. Despite the absence of these parametric knobs, our framework is still capable of creating edge cases that undermine the system’s performance, demonstrating its robustness in edge case generation beyond the reliance on certain powerful and effective parametric knobs (See banner figure on the first page for reference).

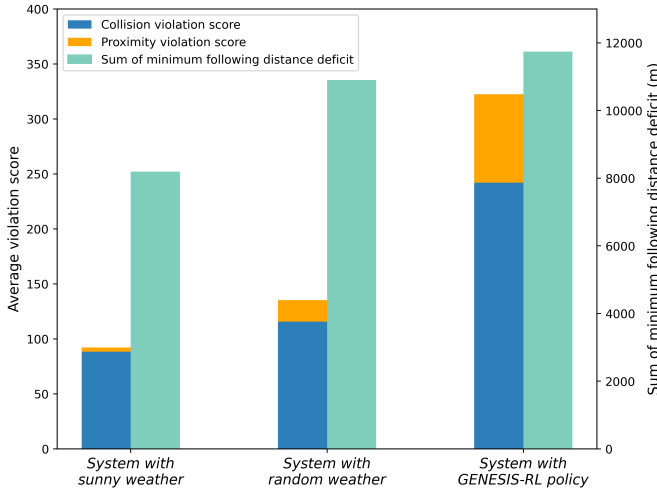


Fig. 2. Violation scores and the sum of minimum following distance deficit (based on RSS) across three testing scenarios - the system operates under sunny weather, random weather and under the GENESIS-RL policy. The results presented are averages across 50 runs with randomly selected initial scenes.

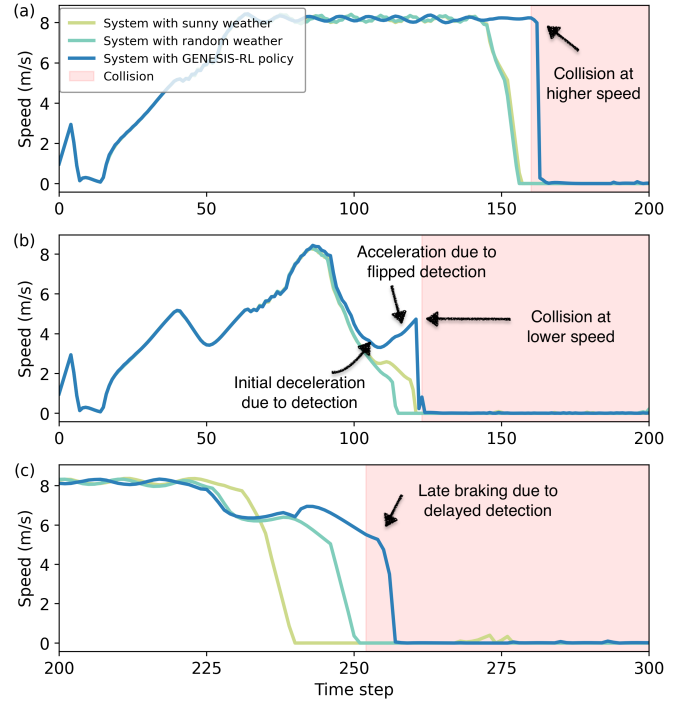


Fig. 3. Examples of the system failure modes based on vehicle telemetry. (a) Example of a non-detection collision - where the system crashed into the vehicle in front at full/high speed due to the non-detection of the other vehicle. (b) Intermittent detection collision - where the intermittent detection prevents the ego vehicle from stopping in time, leading to a lower-speed collision. (c) Delayed detection collision - where the system detects the lead car too late to stop in time. In the tests depicted in this figure, the ego vehicle successfully avoided collisions across both sunny and random weather scenarios and only failed under the conditions generated by GENESIS-RL.

VII. CONCLUSION AND FUTURE WORK

In conclusion, our study demonstrates the GENESIS-RL framework’s capability to generate complex and challenging edge cases for autonomous systems, which are critical for thoroughly testing and enhancing the reliability of systems as such. Moreover, GENESIS-RL proved capable of producing edge cases that will lead to system failure even in the absence of some dominating factors, underscoring its robustness and potential broader application in safety-critical testing environments.

The implications of the results from our work is manifold. First, they highlight the need for including a wide range of challenging scenarios in the testing protocols for autonomous systems, ensuring they are well-equipped to handle the intricacies of complex and dynamic environments. Additionally, the ability of GENESIS-RL to generate test conditions without relying exclusively on dominant factors showcases its value in crafting safer, more dependable autonomous system solutions.

In future work, we aim to expand our exploration to include a wider spectrum of parameters, such as a more comprehensive range of weather parameters, behaviors of other actors (e.g., vehicles and pedestrians) to further enhance the scenario generation capabilities of GENESIS-RL, thereby facilitating the

creation of more diverse and challenging testing environments for autonomous systems.

REFERENCES

- [1] Scott Drew Pendleton, Hans Andersen, Xinxin Du, Xiaotong Shen, Malika Meghiani, You Hong Eng, Daniela Rus, and Marcelo H. Ang. Perception, planning, control, and coordination for autonomous vehicles. *Machines*, 5(1), 2017.
- [2] Dhanoop Karunakaran, Julie Stephany Berrio Perez, and Stewart Worrall. Generating edge cases for testing autonomous vehicles using real-world data. *Sensors*, 24(1):108, 2023.
- [3] Haizhong Zheng, Ziqi Zhang, Juncheng Gu, Honglak Lee, and Atul Prakash. Efficient adversarial training with transferable adversarial examples. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2020.
- [4] Jingkang Wang, Ava Pun, James Tu, Sivabalan Manivasagam, Abbas Sadat, Sergio Casas, Mengye Ren, and Raquel Urtasun. Advsim: Generating safety-critical scenarios for self-driving vehicles. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9909–9918, 2021.
- [5] Yao Deng, Xi Zheng, Tianyi Zhang, Chen Chen, Guannan Lou, and Miryung Kim. An analysis of adversarial attacks and defenses on autonomous driving models. In *2020 IEEE International Conference on Pervasive Computing and Communications (PerCom)*, pages 1–10, 2020.
- [6] Hyung-Jin Yoon, Hamidreza Jafarnejadsani, and Petros Voulgaris. Learning when to use adaptive adversarial image perturbations against autonomous vehicles. *IEEE Robotics and Automation Letters*, 2023.
- [7] Ameya Joshi, Amitangshu Mukherjee, Soumik Sarkar, and Chinmay Hegde. Semantic adversarial attacks: Parametric transformations that fool deep classifiers. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 4773–4783, 2019.
- [8] Augustus Odena, Christopher Olah, and Jonathon Shlens. Conditional image synthesis with auxiliary classifier GANs. In Doina Precup and Yee Whye Teh, editors, *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pages 2642–2651. PMLR, 06–11 Aug 2017.
- [9] Xinquan Chen, Xitong Gao, Juanjuan Zhao, Kejiang Ye, and Cheng-Zhong Xu. Advdiffuser: Natural adversarial example synthesis with diffusion models. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4562–4572, 2023.
- [10] Aditya Ramesh, Mikhail Pavlov, Gabriel Goh, Scott Gray, Chelsea Voss, Alec Radford, Mark Chen, and Ilya Sutskever. Zero-shot text-to-image generation. In Marina Meila and Tong Zhang, editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 8821–8831. PMLR, 18–24 Jul 2021.
- [11] Ming Ding, Zhuoyi Yang, Wenyi Hong, Wendi Zheng, Chang Zhou, Da Yin, Junyang Lin, Xu Zou, Zhou Shao, Hongxia Yang, et al. Cogview: Mastering text-to-image generation via transformers. *Advances in Neural Information Processing Systems*, 34:19822–19835, 2021.
- [12] Jongmin Yoon, Sung Ju Hwang, and Juho Lee. Adversarial purification with score-based generative models. In *International Conference on Machine Learning*, pages 12062–12072. PMLR, 2021.
- [13] Andrew Brock, Jeff Donahue, and Karen Simonyan. Large scale GAN training for high fidelity natural image synthesis. In *International Conference on Learning Representations*, 2019.
- [14] Prafulla Dhariwal and Alexander Nichol. Diffusion models beat gans on image synthesis. *Advances in neural information processing systems*, 34:8780–8794, 2021.
- [15] Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. In *International Conference on Learning Representations*, 2021.
- [16] Qingsheng Zhang and Yongxin Chen. Fast sampling of diffusion models with exponential integrator. In *The Eleventh International Conference on Learning Representations*, 2023.
- [17] Andrea Censi, Konstantin Slutsky, Tichakorn Wongpiromsarn, Dmitry Yershov, Scott Pendleton, James Fu, and Emilio Frazzoli. Liability, ethics, and culture-aware behavior specification using rulebooks. In *2019 International Conference on Robotics and Automation (ICRA)*, pages 8536–8542, 2019.
- [18] Alexey Dosovitskiy, German Ros, Felipe Codevilla, Antonio Lopez, and Vladlen Koltun. CARLA: An Open Urban Driving Simulator. In *Proceedings of the 1st Annual Conference on Robot Learning*, pages 1–16. PMLR, October 2017.
- [19] Baiming Chen, Xiang Chen, Qiong Wu, and Liang Li. Adversarial evaluation of autonomous vehicles in lane-change scenarios. *IEEE Transactions on Intelligent Transportation Systems*, 23(8):10333–10342, 2021.
- [20] Wenhao Ding, Baiming Chen, Minjun Xu, and Ding Zhao. Learning to collide: An adaptive safety-critical scenarios generating method. In *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2243–2250. IEEE, 2020.
- [21] Wenhao Ding, Baiming Chen, Bo Li, Kim Ji Eun, and Ding Zhao. Multimodal safety-critical scenarios generation for decision-making algorithms evaluation. *IEEE Robotics and Automation Letters*, 6(2):1551–1558, 2021.
- [22] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. In *International Conference on Learning Representations*, 2018.
- [23] Amitangshu Mukherjee, Ameya Joshi, Anuj Sharma, Chinmay Hegde, and Soumik Sarkar. Generative semantic domain adaptation for perception in autonomous driving. *Journal of big data analytics in transportation*, 4(2):103–117, 2022.
- [24] Zhang Xinxin, Li Fei, and Wu Xiangbin. Csg: Critical scenario generation from real traffic accidents. In *2020 IEEE Intelligent Vehicles Symposium (IV)*, pages 1330–1336. IEEE, 2020.
- [25] Tommaso Dreossi, Daniel J. Fremont, Shromona Ghosh, Edward Kim, Hadi Ravanbakhsh, Marcell Vazquez-Chanlatte, and Sanjit A. Seshia. VerifAI: A toolkit for the formal design and analysis of artificial intelligence-based systems. In *31st International Conference on Computer Aided Verification (CAV)*, 2019.
- [26] Fabio Muratore, Christian Eilers, Michael Gienger, and Jan Peters. Data-efficient domain randomization with bayesian optimization. *IEEE Robotics and Automation Letters*, 6(2):911–918, 2021.
- [27] Haoyi Niu, Jianming Hu, Zheyu Cui, and Yi Zhang. Dr2l: Surfacing corner cases to robustify autonomous driving via domain randomization reinforcement learning. In *Proceedings of the 5th International Conference on Computer Science and Application Engineering*, pages 1–8, 2021.
- [28] Yuqing Du, Olivia Watkins, Trevor Darrell, Pieter Abbeel, and Deepak Pathak. Auto-tuned sim-to-real transfer. In *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pages 1290–1296. IEEE, 2021.
- [29] Lukas Birkemeyer, Christian King, and Ina Schaefer. Is scenario generation ready for sotif? a systematic literature review. In *2023 IEEE 26th International Conference on Intelligent Transportation Systems (ITSC)*, pages 472–479, 2023.
- [30] Volodymyr Mnih, Koray Kavukcuoglu, David Silver, Alex Graves, Ioannis Antonoglou, Daan Wierstra, and Martin Riedmiller. Playing atari with deep reinforcement learning, 2013.
- [31] Richard S. Sutton and Andrew G. Barto. *Reinforcement Learning: An Introduction*. A Bradford Book, Cambridge, MA, USA, 2018.
- [32] Daniel J. Fremont, Tommaso Dreossi, Shromona Ghosh, Xiangyu Yue, Alberto L. Sangiovanni-Vincentelli, and Sanjit A. Seshia. Scenic: a language for scenario specification and scene generation. In *Proceedings of the 40th ACM SIGPLAN Conference on Programming Language Design and Implementation, PLDI 2019*, page 63–78, New York, NY, USA, 2019. Association for Computing Machinery.
- [33] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms, 2017.
- [34] Antonin Raffin, Ashley Hill, Adam Gleave, Anssi Kanervisto, Maximilian Ernestus, and Noah Dormann. Stable-baselines3: Reliable reinforcement learning implementations. *Journal of Machine Learning Research*, 22(268):1–8, 2021.
- [35] Glenn Jocher. Ultralytics yolov5, 2020.
- [36] Ichiro Hasuo. Responsibility-sensitive safety: an introduction with an eye to logical foundations and formalization. *arXiv preprint arXiv:2206.03418*, 2022.