

TinyLidarNet: 2D LiDAR-based End-to-End Deep Learning Model for F1TENTH Autonomous Racing

Mohammed Misbah Zarrar, Qitao Weng, Bakhbyergyen Yerjan, Ahmet Soyyigit, and Heechul Yun

University of Kansas, Lawrence, KS

{zarrar_1607, wengqt, yerjanb, ahmet.soyyigit, heechul.yun}@ku.edu

Abstract—Prior research has demonstrated the effectiveness of end-to-end deep learning for robotic navigation, where the control signals are directly derived from raw sensory data. However, the majority of existing end-to-end navigation solutions are predominantly camera-based. In this paper, we introduce TinyLidarNet, a lightweight 2D LiDAR-based end-to-end deep learning model for autonomous racing. An F1TENTH vehicle using TinyLidarNet won 3rd place in the 12th F1TENTH Autonomous Grand Prix competition, demonstrating its competitive performance. We systematically analyze its performance on untrained tracks and computing requirements for real-time processing. We find that TinyLidarNet’s 1D Convolutional Neural Network (CNN) based architecture significantly outperforms widely used Multi-Layer Perceptron (MLP) based architecture. In addition, we show that it can be processed in real-time on low-end micro-controller units (MCUs).

I. INTRODUCTION

In F1TENTH autonomous racing [1], [2], developing an intelligent, yet computationally efficient control algorithm is necessary and challenging due to constraints in size, weight, and power. These systems must swiftly process sensor input data to make real-time decisions and allow fast collision-free navigation of the ego-vehicle in various racing tracks. Traditional approaches, which involve a complex pipeline of perception, planning, and control algorithms, are challenging to apply in fast-paced autonomous racing due to high computational costs and sensitivity to perception and modeling errors, which can propagate and accumulate through the pipeline. End-to-end (E2E) deep learning approaches [3]–[6] present a promising alternative as they can simplify the control pipelines, improve computational efficiency, and achieve high performance [6].

In end-to-end approaches, a neural network replaces all or a subset of perception, planning, and control algorithms in the traditional control pipeline [3], [7]. Many prior work has explored end-to-end approaches for autonomous driving [4], [8], [9]. However, the majority of these prior works are vision-based approaches, which require a large amount of training data to achieve good performance and consistency and are susceptible to environmental factors such as lighting conditions [4], [8]. Several studies have investigated end-to-end approaches based on 2D LiDAR (Light Detection and Ranging) in the context of F1TENTH racing specifically [6], [10], [11]. Compared to cameras, 2D LiDARs generate fewer data, making the training of 2D LiDAR models easier.

Most prior 2D LiDAR-based end-to-end approaches for F1TENTH racing are based on multi-layer perceptron (MLP)

models that take a 2D LiDAR scan (e.g., a 1081 dimensional vector over a 270-degree field of view) as input and predict control commands (e.g., throttle and steering) as output [12]. However, these MLP-based models do not perform well at high speeds and do not generalize well across different environments [12].

In this work, we aim to address the following research questions: (1) Can we develop a 2D LiDAR-based end-to-end deep learning model that demonstrates competitive performance in actual F1TENTH autonomous racing competitions? (2) What level of computational power is needed to execute such a deep learning model and achieve competitive performance in racing? (3) Can we achieve good performance on unseen tracks, both in the real world and in simulation, without additional data collection and retraining? In other words, how well does our end-to-end model generalize, especially compared to prior MLP-based models?

To this end, first, we present TinyLidarNet, an end-to-end deep convolutional neural network (CNN) that directly takes a raw 2D LiDAR scan as input and predicts throttle and steering control commands. TinyLidarNet is inspired by the PilotNet architecture [4], a vision-based end-to-end CNN model used in NVIDIA’s Dave 2 project to drive a real car in various real road conditions. Instead of using 2D convolutions, as in the original PilotNet, TinyLidarNet uses 1D convolutions to capture the spatial features of the incoming 2D LiDAR scans. TinyLidarNet shows competitive performance in an actual F1TENTH competition ¹ with a relatively small amount of training data collected on the competition track (Figure 1), following the standard behavior cloning approach [3]. Second, we find that TinyLidarNet is computationally efficient and can perform an inference in less than 1 millisecond on the NVIDIA Jetson NX platform. We further find that it is even possible to execute the network on a tiny MCU in real-time. Specifically, we port the TinyLidarNet to an ESP32-S3 and Raspberry Pi Pico MCU and can achieve sub-50ms latencies (>20Hz) after applying a standard 8-bit integer quantization using TensorFlow Lite-Micro [13]. Lastly, we show that the trained model generalizes well on unseen tracks, both in the real world as well as in simulation, even without any additional rounds of data collection and augmentation regimens such as DAGger [14]. This is in stark contrast with prior MLP-based 2D LiDAR models [6], [10], [11], which struggle to work

¹3rd place winner at 12th F1TENTH autonomous grand-prix competition

well at high speed and unseen tracks [12]. This is because TinyLidarNet’s 1D convolutional filters can better capture the spatial features of the 2D LiDAR scans.

In summary, we make the following contributions:

- We present TinyLidarNet, a 2D LiDAR-based end-to-end CNN architecture for F1TENTH autonomous racing, which shows competitive performance in real racing and generalizes well on unseen tracks. We release the code and training data as open source ².
- We systematically explore the latency, accuracy, and performance trade-offs of TinyLidarNet on contemporary embedded computing platforms. We show the feasibility of using a tiny MCU to run TinyLidarNet in real-time.
- We provide extensive evaluation results of TinyLidarNet’s performance and generalizability, compared to the state-of-the-art 2D LiDAR-based end-to-end MLP models. We show that TinyLidarNet’s CNN architecture is superior to conventional MLP architectures for processing 2D planner LiDAR input.

II. BACKGROUND AND RELATED WORK

A. F1TENTH Autonomous Racing Competition

The F1TENTH Autonomous Racing competition [1] is a competitive autonomous racing event to develop and test algorithms for autonomous racing with 1/10th scale race cars. It challenges the scaled race cars to autonomously navigate and complete races in a given race track as fast as possible without collision. These race cars come equipped with an array of sensors, but the 2D planner LiDAR is the most commonly used primary sensor for perception. The competition provides a realistic and dynamic testing ground for researchers and students in the field of autonomous vehicles, robotics, and artificial intelligence. Figure 1 shows the race track used in the 12th F1TENTH Grand Prix competition held in IEEE/ACM CPS-IoT Week 2023 ³.



Fig. 1: 12th F1TENTH Grand Prix Competition Track

²<https://github.com/CSL-KU/TinyLidarNet>

³<https://cps2023-race.f1tenth.org>

B. End-to-end Deep Learning for Autonomous Driving

Rather than relying on conventional modular systems, end-to-end approaches utilize deep neural networks, transforming sensor input data into control outputs [7]. The concept was initially introduced in the 1980s [15] with a compact 3-layer fully connected neural network tailored for autonomous cars. Subsequent to that, the DARPA Autonomous Vehicle (DAVE) project emerged in the early 2000s [16] and numerous other projects followed the suit [4], [17]–[19]. An example is NVIDIA’s DAVE-2 project, featuring a 9-layer CNN for end-to-end control, demonstrating its ability to drive a real car in various road conditions [4]. Most prior works are vision-based, taking raw image pixels as input and directly generating control commands as output. In [5], on the other hand, a 3D LiDAR is used as the primary input sensor, which produces a 3D point cloud as input to the end-to-end model, in conjunction with GPS and coarse-grain map information. Their end-to-end model directly processes the 3D point cloud via 3D sparse convolutions and was shown to be able to produce robust steering controls in a real autonomous car.

Several prior works have investigated 2D planner LiDAR-based end-to-end approaches in the context of F1TENTH racing [6], [10]–[12]. Compared to cameras or 3D LiDARs, the 2D planner LiDAR in F1TENTH racing cars (see Section III) generates significantly fewer amount data, mainly the distances from the ego vehicle to the surrounding objects or walls on the 2D plane, which makes it easier to process. These prior works targeting F1TENTH racing are based on MLP models to process 2D LiDAR scans. However, in the F1TENTH community, end-to-end approaches have not been popular in actual competitions due to poor performance. For example, in the 12th F1TENTH Grand Prix, we were the only team that adopted an end-to-end approach. In [20], it is shown that a CNN model can be trained more efficiently than an MLP model in a deep reinforcement learning framework, but it did not show high-speed racing capability in the real world. In a recent survey by the creators of the F1TENTH racing [7], it is noted that end-to-end approaches exhibit low achievable speed, non-ideal driving characteristics (e.g., wobbly behavior), and low generalizability on unseen tracks, among other limitations. In this work, we challenge these preconceived weaknesses of end-to-end approaches in the context of F1TENTH Racing.

III. F1TENTH PLATFORM AND 2D LiDAR

Figure 2 shows the F1TENTH car utilized in all of our real-world experiments. The chassis of the car is based on a Traxxas Rally 1/10-scale radio-controlled car with an Ackermann steering mechanism. Its primary sensor is a Hokuyo UST-10LX 2D Planner LiDAR, which has a 270° field of view with a range of up to 10 meters. This LiDAR device provides an angular resolution of 0.25° and operates at a scan frequency of 40Hz, generating 1081 data points represented as a 1D distance array, which serves as the input for our end-to-end model. Since the scan frequency is 40 Hz, naturally the deadline for prediction of speed and throttle becomes 0.025 seconds or 25 milliseconds. For control, an open-source electronic speed control (ESC) board controls

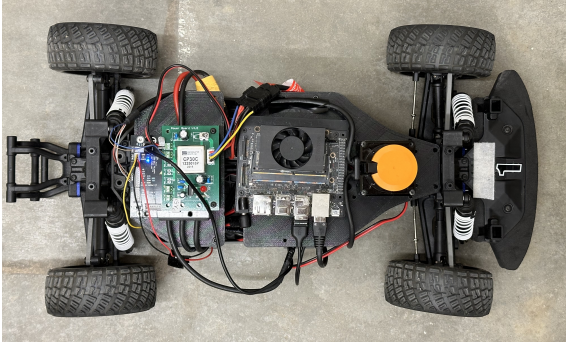


Fig. 2: F1TENTH platform with a Hokuyo UST-10LX 2D LiDAR and a NVIDIA Jetson Xavier NX on-board computer.

the RPM of a brushless motor and a steering servo. A power board distributes power from a lithium polymer (LiPo) battery to the sensors, motors, and the on-board computer. For on-board computing, an NVIDIA Jetson Xavier NX embedded computer is used to run all software, which equips 8 64-bit ARM CPU cores and a 384-core Volta GPU. On the software side, we use the ROS Noetic Ninjemys framework on Ubuntu 20.04.6 operating system.

IV. TINYLIDARNET

In this section, we introduce the TinyLidarNet architecture and our data collection and training process.

A. Architecture

Figure 3 shows the architecture of TinyLidarNet. The architecture is comprised of 9 layers (5 convolutional and 4 fully connected) with a total of 220K parameters and takes about 1.5 million multiply-accumulate operations (MACs) to execute. The architecture is inspired by the PilotNet architecture [4], which is a vision-based end-to-end model for NVIDIA’s Dave2 self-driving car. The main differences are that it takes a 2D LiDAR scan, instead of an RGB image, as input and uses 1D CNN layers instead of 2D ones. Quantitatively, compared to PilotNet [4], TinyLidarNet has a similar number of parameters (250K vs. 220K) but requires significantly lower MACs (27 million vs. 1.5 million) because the input dimensionality is much lower (200x66 RGB pixels vs. 1081 LiDAR scan samples). As such, it has a significantly lower computational demand, which makes it possible to use less powerful computing platforms, as we investigate in Section V-C.

B. Data Collection, Pre-processing and Training

The data collection is carried out by driving the vehicle manually using a joystick. The collected data includes servo angle and speed data from the ESC board of the car and the 1D range array (1081 values) from the Hokuyo 2D LiDAR.

We collected a total of 12,329 samples, approximately 5 minutes of driving on the training track, shown in Figure 1. Note that during the data collection, the car drove with stationary obstacles (opponent vehicles) present. We assigned

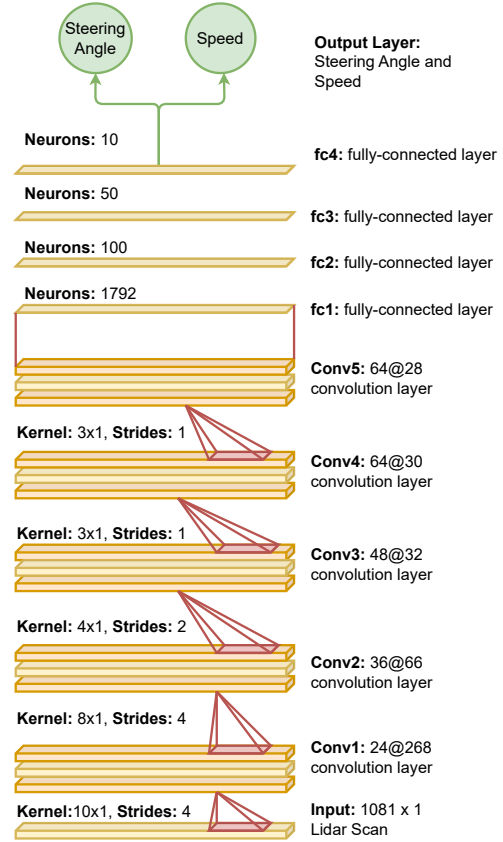


Fig. 3: TinyLidarNet architecture: 9 layers (5 convolutional, 4 fully-connected) with 220,686 parameters.

85% (10,478 samples) of the collected samples to training and the remaining portion to testing and validation. We used a batch size of 64 and trained the network for 20 epochs with a learning rate of $5e-5$ and Huber loss [21].

The TinyLidarNet training process is performed locally on the NVIDIA Jetson Xavier NX platform and takes approximately 4 minutes to complete.

During the data collection process, we found that the quality of the LiDAR scan samples deteriorated significantly when the car navigated dark track sections. LiDAR sensors tend to perform less effectively on black surfaces, resulting in noisy point clouds [22]. To reduce noise, we employ median and interpolation filters. These filters play a crucial role in stabilizing point clouds.

V. EVALUATION

In this section, we evaluate the performance of TinyLidarNet.

A. Insights from an F1TENTH Competition

Using TinyLidarNet, we participated in the 12th F1TENTH Autonomous Grand Prix competition and won the 3rd place award out of 13 participating teams. Although TinyLidarNet did not secure the top position, it exhibited competitive performance and demonstrated several advantages over traditional approaches that relied on perception, planning, and control pipelines.

First, during the competition, the track depicted in Figure 1 underwent frequent alterations due to collisions. Each time a racecar collided with the track, the track was manually adjusted, leading to slight changes in its configuration. This posed a considerable challenge for participants who relied on an offline map for path planning. In contrast, TinyLidarNet does not rely on an offline map and was largely immune to the alterations in the track configuration.

Second, a particularly interesting capability of TinyLidarNet is its ability to overtake other cars during head-to-head races in the tournament. This is interesting because the training dataset comprises examples of the car driving on the competition track by a human driver with static obstacles only, as detailed in Section IV-B. Note that executing overtaking maneuvers in high-speed racing is a challenging task for classical control approaches [23], [24], and most participants in the competition did not attempt to execute overtaking maneuvers during the race. A team that implemented an overtaking algorithm did so only in the long straight section of the track (Figure 1). On the other hand, TinyLidarNet was able to overtake moving opponents in any part of the track throughout the competition, likely because moving race cars were recognized as static obstacles or walls.

Overall, the competitive performance of TinyLidarNet in the competition motivated us to systematically analyze its performance and characteristics.

B. Performance on Simulated Tracks

In this subsection, we evaluate TinyLidarNet’s performance on unseen virtual tracks.

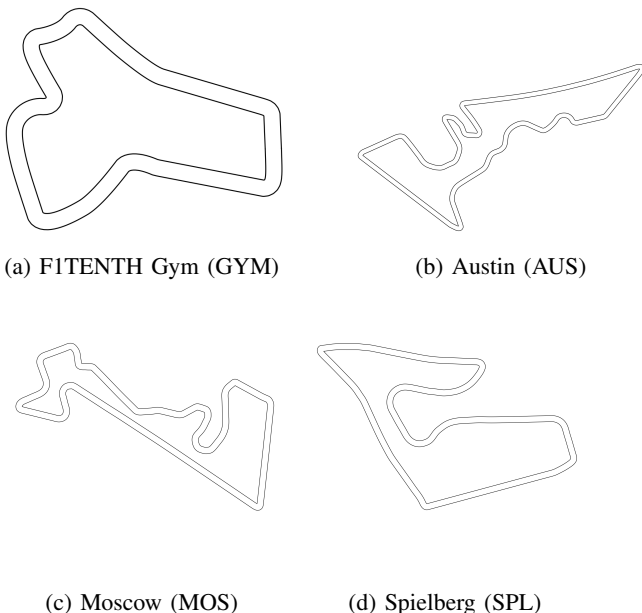


Fig. 4: Simulation Tracks from F1TENTH gym and F1TENTH racetrack repository [25], [26]

Table I shows the basic characteristics of the compared 2D Lidar-based end-to-end models. Note that *TinyLidarNet^L* is the original version that we use in an actual competition V-A. It is also the largest one as it takes all 1081 range data of a

Model	Input dim.	Parameters	MACs
TinyLidarNet ^L	1081	220,686	1,546,960
TinyLidarNet ^M	541	111,886	687,680
TinyLidarNet ^S	271	54,286	240,752
MLP256 ^L [12]	1081	343,298	342,784
MLP256 ^M	541	205,058	204,544
MLP256 ^S	271	135,938	135,424

TABLE I: Comparison of end-to-end models.

LiDAR scan as input. *TinyLidarNet^M* and *TinyLidarNet^S* are smaller variants, which down-sample the LiDAR scan by taking one for every 2 or 4 values of the range data, respectively. Except for the input dimension, they have identical 9-layer architecture. On the other hand, *MLP256^L* [12] is a multi-layer perceptron with 2 hidden layers, with 256 neurons in each layer. *MLP256^M* and *MLP256^S* are its smaller variants with smaller input dimensions.

Figure 4 shows the four tracks we use for evaluation. These tracks are widely used in F1TENTH literature [25], [26], which we use without any modifications.

The basic evaluation setup is as follows. First, we train all compared models using the same dataset, which was originally collected on the F1TENTH competition, as described in Section IV-B, and evaluate them on the four tracks using a simulator infrastructure in [27]. Using the simulation infrastructure, each end-to-end model drives a simulated car for one complete lap on a track for 10 trials, each time starting from a random position on the track. For each successful completion among the 10 trials, the *average lap time* was calculated. In addition, the *average progress rate* measures the percentage of the track the model can progress, on average, during the 10 trials.

To mimic real-world conditions, we adjusted the simulation so that the control frequency is set to 40 Hz. We also add random Gaussian noise in the range of 0 to 0.5 to the LiDAR scans and perform a range clipping at 10 m to mimic the behavior of Hokuyo UST-10LX LiDAR in the F1TENTH platform, which has a range of 10 m.

Note that a common practice of training an end-to-end model for a robot involves starting in simulation and eventually deploying it in the real world. However, bridging the Simulation to Reality gap (Sim2Real gap) can pose significant challenges [28], [29]. On the other hand, our evaluation strategy can be described as Real2Sim as we train the models using the real-world track dataset and test them in a simulated environment.

Table II shows the evaluation results. Note first that all three *TinyLidarNet^{L,M,S}* can complete all 10 laps without any collision, hence achieving 100% average progress, in all four tested tracks. On the other hand, MLP256 models struggle to complete the laps. MLP256 models could complete the lap some of the times, but their success rates are not close to 100% in most cases. The average lap times of MLP256 models are also somewhat slower than that of TinyLidarNet models in many cases. Overall, the TinyLidarNet variants show better performance, especially in terms of average progress, than the MLP256 variants.

Model	Average Lap Time (s)				Average Progress (%)			
	GYM	AUS	MOS	SPL	GYM	AUS	MOS	SPL
TinyLidarNet ^L	25.8	85.7	63.3	65.3	100	100	100	100
TinyLidarNet ^M	25.3	80.0	59.5	61.5	100	100	100	100
TinyLidarNet ^S	26.9	83.4	61.8	64.1	100	100	100	100
MLP256 ^L [12]	N/A	N/A	58.8	58.3	31	16	42	61
MLP256 ^M	28.4	N/A	64.3	65.7	100	17	58	78
MLP256 ^S	27.6	N/A	N/A	62.2	77	48	29	37

TABLE II: Average lap time (of successful trials) and the average progress (of all trials). At each trial, the ego-vehicle’s position in the track is randomly chosen.

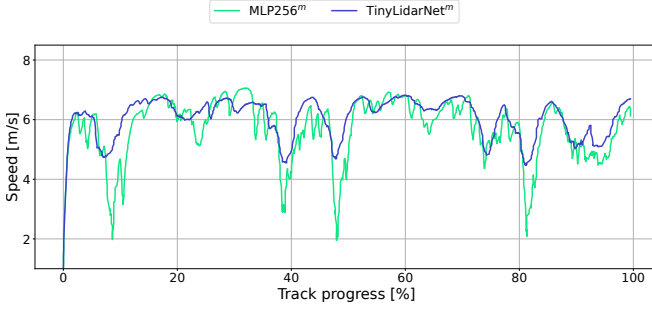


Fig. 5: Speed comparison of different models on GYM Track

The Figure 5 and Figure 6 show the driving speed and the trajectory, respectively, of the TinyLidarNet^M and MLP256^M models in one of the successfully completed laps out of 10 trials on the GYM track. Note that MLP256^M’s speed fluctuates significantly and its trajectory is wobbly. In contrast, TinyLidarNet^M maintains a more consistent speed and stable trajectory. Note that in terms of average progress, all TinyLidarNet models consistently outperform MLP256 models in all evaluated tracks. Regarding average lap time, TinyLidarNet models are usually faster than or similar to MLP256 models.

In summary, we show that TinyLidarNet models, trained on real-world data, are well generalized to navigate unseen race tracks, while the commonly used MLP models struggle to generalize. This is because the CNN layers of TinyLidarNet can capture spatial features of LiDAR scans, enabling better generalization.

C. Inference Latency

In the context of racing, latency is of paramount importance due to the inherent stringent real-time requirements in such applications. In this subsection, we evaluate the inference latency of TinyLidarNet models on different computing platforms.

	Xavier NX	ESP32-S3	RPi Pico
CPU	NVIDIA Carmel 6C@1.9 GHz	Xtensa LX7 2C@240 MHz	ARM Cortex-M0+ 2C@133 Mhz
Memory	8GB LPDDR4x	8MB PSRAM	264KB SRAM
Storage	16GB eMMC	8MB Flash	2MB Flash

TABLE III: Computing platforms

Table III shows the three computing platforms we use for evaluation. Note that Jetson Xavier NX is the main onboard

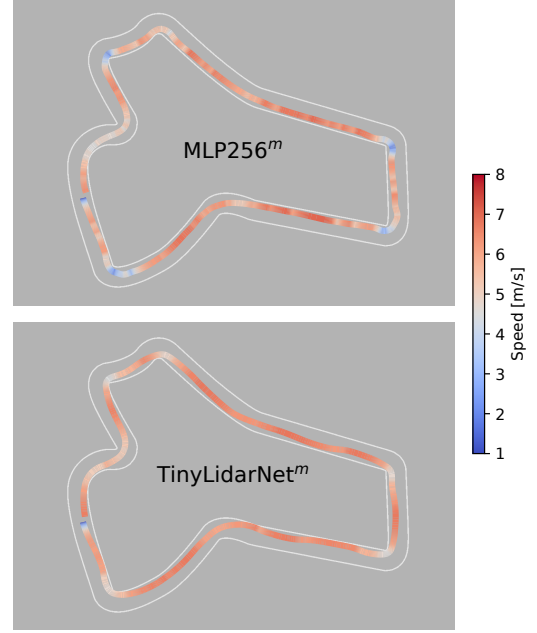


Fig. 6: A snippet of the trajectory of different models with speed profile on GYM track.

computer platform of our F1TENTH platform, which features powerful six ARM CPU cores running at 1.9 GHz. We also use XIAO ESP32-S3 [30] and Raspberry Pi Pico [31] MCUs to understand the potential of using these inexpensive low-end computing platforms for the real-time processing of TinyLidarNet. To this end, we port all three versions of TinyLidarNet on the MCUs and measure the inference latency of the models.

Model	Xavier NX	ESP32-S3	RPi Pico
TinyLidarNet ^L (fp32)	<1	838	2642
TinyLidarNet ^L (int8)	<1	16	196
TinyLidarNet ^M (int8)	<1	8	91
TinyLidarNet ^S (int8)	<1	4	36

TABLE IV: Inference latency (ms) comparison

Table IV shows the results. Note that (fp32) refers to the baseline 32-bit floating point number-based models whereas (int8) refers to 8-bit integer quantized models to reduce computational overhead and storage space demand to store the weights of the model.

First, note that all TinyLidarNet models are small enough to fit in both MCUs. Even in the smaller Raspberry Pi Pico MCU with only 264KB SRAM and 2MB Flash, all TinyLidarNet models can fit without any issue.

Second, the inference latency of TinyLidarNet^L (fp32), however, is more than 2 seconds on the Pico and more than 800 ms on the ESP32-S3 MCU, which is more powerful than Pico but still significantly limited compared to Xavier NX.

Third, using 8-bit quantized models enables TinyLidarNet to be executed in real-time as the TinyLidarNet^L (int8) model takes only 16 ms to perform an inference on the ESP32-S3. This means that the model can run at >50Hz. The inference latency of the smallest TinyLidarNet^S (int8) model is only 36 ms (>20Hz) on the Pico, which is sufficient for real-time control for most robotics applications.

In summary, with a standard quantization-based optimization, we show that TinyLidarNet is lightweight enough to run on low-end MCUs in real-time.

D. Performance on Unseen Real Tracks

In this subsection, we evaluate the performance of TinyLidarNet on an unseen real track to validate the findings on the simulated tracks in Section V-B.

As in Section V-B, the evaluated models are all trained using the dataset we collected during the F1TENTH competition, described in Section IV-B. The trained models are then tested on a different track installed in our lab. Figure 1 shows the track.



Fig. 7: Real-world test track

We employed a similar performance metric to what was discussed in Section V-B, conducting these experiments over 5 trials with the ego vehicle placed on the track at 5 randomly selected starting positions, indicated by the red markers visible in Figure 7. To ensure a fair comparison, we set the minimum speed to -0.5 m/s, considering the vehicle’s inertia, and capped the maximum speed limit at 5 m/s, aligning with the linearly mapped data range from 0 m/s to 5 m/s used during training for all models.

Table V shows the results. Across five trials, all MLP256 variants consistently crashes, highlighting the struggle to generalize on a new track that was not seen during training. All three MLP256 models were unable to complete a single lap and struggled especially when they needed to take hard

Model	Average Lap Time (s)	Success Rate (%)
TinyLidarNet ^L	20.5	100
TinyLidarNet ^M	19.9	100
TinyLidarNet ^S	19.5	80
MLP256 ^L [12]	N/A	0
MLP256 ^M	N/A	0
MLP256 ^S	N/A	0

TABLE V: Average lap time and success rate of the models on the real track. The average lap time was calculated from the successful lap completion of 5 trials by placing the ego vehicle in random positions on the map.

U-shaped turns at the bottom left and top right corners of the track shown in Figure 7.

In contrast, all three versions of TinyLidarNet were able to complete the lap. TinyLidarNet^S crashed only once (out of five trials), possibly because it is the smallest of all and loses 3/4 of the LiDAR range information due to downsampling (skipping three out of every four range values). Nevertheless, all TinyLidarNet models clearly outperform the MLP model.

As discussed earlier, TinyLidarNet’s superior performance can be attributed from its use of 1D CNN layers, which can capture spatial features of 2D LiDAR scans better. As a result, TinyLidarNet models generalize well on unseen tracks as well as unseen moving vehicles, all of which have distinct spatial features that can be learned. While it is well-known that CNN is good at processing vision data, our results strongly suggest that it is also good at processing 2D LiDAR data, which is consistent with recent findings in [20].

VI. CONCLUSION

In this study, we introduced TinyLidarNet, a lightweight 2D LiDAR-based end-to-end deep learning model designed for F1TENTH racing. TinyLidarNet takes an array of LiDAR scans as direct input and predicts steering and speed through a 1D convolutional neural network.

We conducted a comprehensive evaluation on the generalizability of the model on unseen tracks both in simulation and in the real-world. We showed that TinyLidarNet’s 1D CNN-based architecture is superior to the commonly used MLP architecture in processing 2D LiDAR scan data.

We also evaluated the effects of quantization on TinyLidarNet. Deploying it on two low-cost MCUs, we found its inference to be sufficiently efficient for autonomous vehicle control, demonstrating its lightweight nature and low latency.

Future research avenues include exploring the use of DAGger [14], [32] to improve data collection, and using TinyLidarNet as a foundation for bootstrapping with Deep Reinforcement Learning techniques, further enhancing its capabilities and adaptability in challenging racing scenarios.

ACKNOWLEDGMENTS

This research is supported in part by the NSF grant CPS-2038923.

REFERENCES

- [1] F1Tenth, “F1/10 autonomous racing competition.” <http://f1tenth.org>.
- [2] M. O’Kelly, H. Zheng, D. Karthik, and R. Mangharam, “F1Tenth: An open-source evaluation environment for continuous control and reinforcement learning,” *Proceedings of Machine Learning Research*, vol. 123, 2020.
- [3] S. Levine, C. Finn, T. Darrell, and P. Abbeel, “End-to-end training of deep visuomotor policies,” *Journal of Machine Learning Research*, 2016.
- [4] M. Bojarski *et al.*, “End-to-End Learning for Self-Driving Cars,” *arXiv*, 2016.
- [5] Z. Liu, A. Amini, S. Zhu, S. Karaman, S. Han, and D. L. Rus, “Efficient and robust lidar-based end-to-end navigation,” in *2021 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 13247–13254, IEEE, 2021.
- [6] L. Chen, P. Wu, K. Chitta, B. Jaeger, A. Geiger, and H. Li, “End-to-end autonomous driving: Challenges and frontiers,” 2023.
- [7] J. Betz, H. Zheng, A. Liniger, U. Rosolia, P. Karle, M. Behl, V. Krovi, and R. Mangharam, “Autonomous vehicles on the edge: A survey on autonomous vehicle racing,” *IEEE Open Journal of Intelligent Transportation Systems*, vol. 3, pp. 458–488, 2022.
- [8] M. Bojarski, C. Chen, J. Daw, A. Degirmenci, J. Deri, B. Firner, B. Flepp, S. Gogri, J. Hong, L. Jackel, Z. Jia, B. Lee, B. Liu, F. Liu, U. Muller, S. Payne, N. K. N. Prasad, A. Provodin, J. Roach, T. Rvachov, N. Tadimeti, J. van Engelen, H. Wen, E. Yang, and Z. Yang, “The nvidia pilotnet experiments,” 2020.
- [9] Y. Pan, C.-A. Cheng, K. Saigol, K. Lee, X. Yan, E. Theodorou, and B. Boots, “Agile autonomous driving using end-to-end deep imitation learning,” 2019.
- [10] F. Codevilla, M. Müller, A. López, V. Koltun, and A. Dosovitskiy, “End-to-end driving via conditional imitation learning,” 2018.
- [11] S. N. Wadekar, B. J. Schwartz, S. S. Kannan, M. Mar, R. K. Manna, V. Chellapandi, D. J. Gonzalez, and A. E. Gamal, “Towards end-to-end deep learning for autonomous racing: On data collection and a unified architecture for steering and throttle prediction,” 2021.
- [12] X. Sun, M. Zhou, Z. Zhuang, S. Yang, J. Betz, and R. Mangharam, “A benchmark comparison of imitation learning-based control policies for autonomous racing,” in *2023 IEEE Intelligent Vehicles Symposium (IV)*, pp. 1–5, IEEE, 2023.
- [13] R. David, J. Duke, P. Warden, *et al.*, “Tensorflow lite micro: Embedded machine learning on tinyml systems,” *arXiv preprint arXiv:2010.08678*, 2020.
- [14] S. Ross, G. Gordon, and D. Bagnell, “A reduction of imitation learning and structured prediction to no-regret online learning,” in *Proceedings of the fourteenth international conference on artificial intelligence and statistics*, pp. 627–635, JMLR Workshop and Conference Proceedings, 2011.
- [15] D. a. Pomerleau, “ALVINN: An autonomous land vehicle in a neural network,” in *Advances in Neural Information Processing Systems (NIPS)*, 1989.
- [16] Y. Lecun, E. Cosatto, J. Ben, U. Muller, and B. Flepp, “DAVE: Autonomous off-road vehicle control using end-to-end learning,” Tech. Rep. DARPA-IPTO Final Report, Courant Institute/CBLL, 2004.
- [17] M. G. Bechtel, E. McElhiney, M. Kim, and H. Yun, “Deepcar: A low-cost deep neural network-based autonomous car,” in *IEEE International Conference on Embedded and Real-Time Computing Systems and Applications (RTCSA)*, 2018.
- [18] M. Bechtel, Q. Weng, and H. Yun, “Deepcar: Applying tinyml to autonomous cyber physical systems,” in *2022 IEEE 28th International Conference on Embedded and Real-Time Computing Systems and Applications (RTCSA)*, pp. 120–127, IEEE, 2022.
- [19] T. Weiss and M. Behl, “Deepcar: A framework for autonomous racing,” in *2020 Design, automation & test in Europe conference & exhibition (DATE)*, pp. 1163–1168, IEEE, 2020.
- [20] M. Bosello, R. Tse, and G. Pau, “Train in austria, race in montecarlo: Generalized rl for cross-track f1 tenth lidar-based races,” in *2022 IEEE 19th Annual Consumer Communications & Networking Conference (CCNC)*, pp. 290–298, IEEE, 2022.
- [21] P. J. Huber, “Robust estimation of a location parameter,” *Annals of Mathematical Statistics*, vol. 35, pp. 492–518, 1964.
- [22] S. Wu, G. K. Reddy, and D. Banerjee, “Pitch-black nanostructured copper oxide as an alternative to carbon black for autonomous environments,” *Advanced Intelligent Systems*, vol. 3, no. 9, p. 2100049, 2021.
- [23] S. Bak, J. Betz, A. Chawla, H. Zheng, and R. Mangharam, “Stress testing autonomous racing overtake maneuvers with rrt,” in *2022 IEEE Intelligent Vehicles Symposium (IV)*, pp. 806–812, IEEE, 2022.
- [24] J. Zhang and H.-W. Loidl, “F1Tenth: An over-taking algorithm using machine learning,” in *International Conference on Automation and Computing (ICAC)*, pp. 01–06, IEEE, 2023.
- [25] M. O’Kelly, H. Zheng, D. Karthik, and R. Mangharam, “F1Tenth: An open-source evaluation environment for continuous control and reinforcement learning,” in *NeurIPS 2019 Competition and Demonstration Track*, pp. 77–89, PMLR, 2020.
- [26] J. Betz, H. Zheng, A. Liniger, U. Rosolia, P. Karle, M. Behl, V. Krovi, and R. Mangharam, “Autonomous vehicles on the edge: A survey on autonomous vehicle racing,” *IEEE Open J. Intell. Transp. Syst.*, vol. 3, pp. 458–488, 2022.
- [27] B. D. Evans, R. Trumpp, M. Caccamo, H. W. Jordaan, and H. A. Engelbrecht, “Unifying f1tenth autonomous racing: Survey, methods and benchmarks,” 2024.
- [28] P. Trentsios, M. Wolf, and D. Gerhard, “Overcoming the sim-to-real gap in autonomous robots,” *Procedia CIRP*, vol. 109, pp. 287–292, 2022. 32nd CIRP Design Conference (CIRP Design 2022) - Design in a changing world.
- [29] X. Hu, S. Li, T. Huang, B. Tang, R. Huai, and L. Chen, “How simulation helps autonomous driving: A survey of sim2real, digital twins, and parallel intelligence,” *IEEE Transactions on Intelligent Vehicles*, vol. 9, no. 1, pp. 593–612, 2024.
- [30] “Seed Studio XIAO ESP32S3.” <https://www.seeedstudio.com/XIAO-ESP32S3-p-5627.html>.
- [31] “Raspberry Pi Pico and Pico W.” <https://www.raspberrypi.com/documentation/microcontrollers/raspberry-pi-pico.html>.
- [32] M. Kelly, C. Sidrane, K. Driggs-Campbell, and M. J. Kochenderfer, “Hg-dagger: Interactive imitation learning with human experts,” in *2019 International Conference on Robotics and Automation (ICRA)*, pp. 8077–8083, 2019.