

Localization-Aware Adversarial Attacks Against LiDAR-Based 3D Object Detection

Peng Gao¹, Guanhua Sun²[0009-0004-7186-0662], Shuo Chai¹, and Tieying Zhu^{1,*}

¹ School of Information Science and Technology, Northeast Normal University,
Changchun, China

² Department of Data Science, City University of Hong Kong,
Hong Kong SAR, China

* Corresponding author
zhuty@nenu.edu.cn

Abstract. Robustness in 3D object detection is paramount for safety-critical applications such as autonomous driving. However, existing adversarial attack methodologies often fail to fully leverage the intrinsic spatial geometric characteristics of 3D point clouds. To address this limitation, we propose a novel differentiable method named Localization-Aware Adversarial Attacks (LAA). LAA explicitly incorporates the absolute coordinates of bounding boxes, the relative spatial relationships with ground truth objects, and rotation angles into its adversarial optimization objective. By generating imperceptible point cloud perturbations, LAA aims to directly disrupt the localization awareness of 3D object detectors. Extensive experiments on the KITTI dataset against a variety of mainstream 3D object detectors demonstrate that LAA exhibits remarkable effectiveness in compromising the detectors' localization capabilities. This research reveals the vulnerability of current 3D object detectors regarding localization awareness and provides valuable insights for constructing more robust detection systems.

Keywords: 3D object detection, point cloud, LiDAR, adversarial attacks.

1 Introduction

In recent years, LiDAR-based 3D object detection has emerged as a cornerstone technology for autonomous driving. Driven by the advancements in Deep Neural Networks (DNNs), this technology has achieved remarkable performance in perceiving and localizing surrounding objects from point cloud data [1,2,3]. The establishment of standardized benchmarks, such as KITTI [4] and nuScenes [5], has accelerated progress in this field by providing comprehensive evaluation frameworks for 3D object detectors. However, the inherent vulnerability of DNNs to adversarial examples presents a significant and unresolved security risk [6,7]. These robustness challenges are further compounded by the intrinsic complexity of 3D point cloud processing, which involves irregular data structures and high-dimensional feature

spaces [8]. When deployed in autonomous systems, such vulnerabilities could be exploited to cause catastrophic perception failures, underscoring the urgent need to rigorously evaluate and fortify the robustness of these critical systems [9].

Adversarial attacks have been extensively studied in the 2D image domain [10,11], with pioneering works [12,13] successfully extending these methods to 3D point cloud classification, establishing fundamental attack paradigms [14]. However, adapting these methods to the 3D object detection task presents unique challenges, necessitating improvements tailored to the specific scenes and objectives of detection. Zhang et al. [15] conducted a comprehensive study extending point cloud classification attacks to the detection task. However, their approach primarily inherits the classification attack philosophy to maximize detection loss; the attack on localization awareness remains at the level of naive coordinate regression, lacking deep exploitation of 3D geometric semantics. Chen et al. [16] argued that current attacks overlook the disruption of the detector’s perception capabilities and thus introduced classification confidence and IoU to construct an adversarial loss. However, in 3D space, IoU cannot precisely represent the spatial relationship between the detection box and the ground truth, and it suffers from the gradient vanishing problem when bounding boxes do not overlap.

Currently, most adversarial attack methods rely on aggregate metrics like mean Average Precision (mAP) to evaluate effectiveness. However, such metrics fail to distinguish whether performance degradation stems from classification errors or localization errors. In fact, existing attacks tend to induce a proliferation of low-quality False Positive predictions. While this drastically reduces aggregate metrics like mAP, it often fails to prevent the detector from correctly recalling true objects. Consequently, the detector remains capable of detecting real obstacles, meaning the perception system remains functionally operative. Fig. 1 illustrates the overview of our proposed localization-aware adversarial attack pipeline.

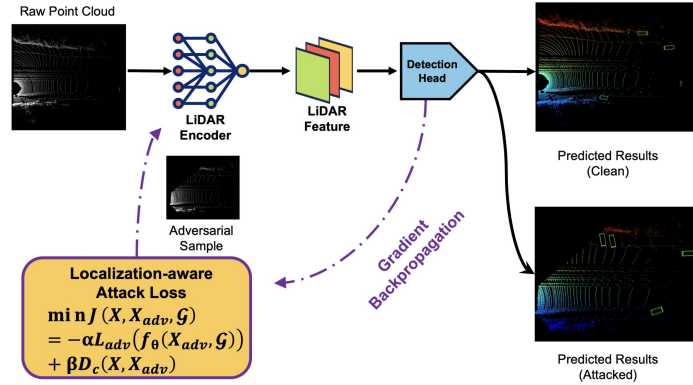


Fig. 1. Overview of the proposed localization-aware adversarial attack pipeline. The solid black line indicates the vanilla 3D detection inference process. The adversarial sample is generated by iteratively updating the input point cloud via gradients backpropagated from our proposed Localization-Aware Adversarial Loss (blue dashed line).

Based on the above analysis, our main contributions are as follows:

- We propose Localization-Aware Adversarial Attack (LAA), as illustrated in Fig. 1. To exploit the geometric properties of 3D point clouds, LAA incorporates bounding box coordinates, spatial relationships with ground truth, and rotation angles into the optimization objective. Through a gradient-based loss function, LAA generates imperceptible perturbations that directly disrupt the model’s ability to localize 3D bounding boxes.
- We evaluate the robustness of 3D detectors across diverse architectures. In addition to mAP, we introduce the Location Attack Success Rate (ASR) to measure localization failures specifically. Our analysis shows that these localization errors, reflected by high ASR, are the main cause of overall detection performance degradation.

Open Science Statement. To support reproducibility and facilitate future research, we will make our source code and experimental configurations publicly available upon acceptance of this paper.

2 Related Work

2.1 LiDAR-Based 3D Object Detection

Current 3D object detectors can be broadly categorized by their point cloud representation strategy, with comprehensive surveys providing detailed taxonomies of these approaches [8,14]. Voxel-based methods first discretize the irregular point cloud into a regular grid or voxel, enabling the efficient application of 3D or 2D convolutions. Seminal works like VoxelNet [3] and the highly efficient PointPillars [1] established this paradigm, while SECOND [17] introduced sparse convolution techniques to improve computational efficiency. Point-based detectors, conversely, operate directly on raw, unstructured points to preserve precise geometric information, as demonstrated by architectures like PointRCNN [18]. To leverage the strengths of both, hybrid methods such as PV-RCNN [2] fuse voxel-based and point-level features to achieve high performance. While these detectors have rapidly advanced in accuracy, the adversarial robustness of these diverse architectures remains insufficiently explored, a gap highlighted by recent studies [15].

2.2 Adversarial Attacks on 3D Object Detection

The study of adversarial attacks originated in the 2D image domain, with foundational works revealing the inherent vulnerabilities of deep neural networks [6,7]. Subsequent research developed sophisticated attack methodologies [10] and evaluation frameworks [11], significantly advancing our understanding of adversarial robustness in computer vision. This line of inquiry was later extended to the 3D domain, where initial efforts targeted the simpler task of point cloud classification and established basic attack paradigms such as point perturbation, detachment, and attachment [12,13].

To disrupt the perception capabilities of detectors, Chen et al. [16] introduced a joint adversarial loss incorporating both classification confidence and the Intersection-over-Union (IoU) between predicted boxes and ground truths. However, the IoU metric struggles to precisely capture complex spatial relationships in 3D space. Furthermore, directly optimizing IoU faces the gradient vanishing problem when the bounding boxes do not overlap, hindering effective perturbation generation. Moreover, research has demonstrated that digital attacks are physically realizable, successfully deceiving detectors in real-world scenarios [19] and posing direct risks to autonomous systems [9], which further exacerbates the severity of this threat. In this work, we focus primarily on LiDAR data to generate adversarial examples against 3D object detection models within the digital domain.

2.3 Limitations of IoU in 3D Space

The Intersection-over-Union (IoU) metric has been widely adopted for bounding box evaluation in both 2D and 3D object detection tasks. However, when applied to adversarial attack optimization in 3D space, IoU exhibits several inherent limitations that hinder its effectiveness as an attack objective:

- (1) **Gradient Vanishing for Non-overlapping Boxes.** When the predicted bounding box and the ground-truth box do not intersect, IoU yields a value of zero with undefined gradients. This renders gradient-based optimization infeasible in scenarios where the adversarial perturbation causes the prediction to deviate significantly from the ground truth. Consequently, attacks relying solely on IoU cannot effectively guide perturbations when boxes become disjoint.
- (2) **Insensitivity to Rotation.** Standard 3D IoU calculation treats the orientation angle as a secondary factor in the overlap computation. Two pairs of bounding boxes may exhibit identical IoU values while having drastically different orientations. For autonomous driving applications, accurate heading estimation is critical for downstream motion planning and trajectory prediction tasks. An attack that induces rotational errors without affecting IoU fails to expose this vulnerability.
- (3) **Scale Ambiguity.** IoU normalizes the intersection by the union volume, making it inherently scale-invariant. However, for safety-critical perception systems, the absolute positional error matters more than the relative overlap. A 1-meter localization offset for a nearby pedestrian poses a significantly greater collision risk than the same offset for a distant truck, yet both scenarios may yield similar IoU degradation.

These limitations motivate our design of the localization-aware loss L_{loc} , which explicitly addresses positional offset through center distance terms (ρ^2/c^2), incorporates rotational deviation via rotation-aware components (ρ_{rot}^2/c_{rot}^2), and captures relative spatial alignment beyond simple overlap metrics.

3 Approach

We begin by formalizing the adversarial attack problem. The input point cloud to a LiDAR-based 3D object detector is denoted as $X = \{p_i \mid i = 1, \dots, N\}$, where we temporarily omit physical attributes such as intensity and elongation. Each point $p_i \in \mathbb{R}^3$ is represented solely by its three-dimensional spatial coordinates (x_p, y_p, z_p) . The ground-truth annotations for the scene consist of a set of 3D bounding boxes $G = \{B_j\}$, where each bounding box $B_j \in \mathbb{R}^8$ is a vector containing the object center coordinates (x_j, y_j, z_j) , dimensions (w_j, l_j, h_j) , and orientation angle r_j .

Our objective is to design an adversarial attack algorithm that generates adversarial samples $X_{adv} = \{X + \Delta X \mid X, \Delta X \in \mathbb{R}^3\}$ by applying small and imperceptible perturbations ΔX to the original point cloud X . These adversarial samples should produce erroneous outputs when fed into the detector $f_\theta(\cdot)$. The specific algorithmic process refers to Algorithm 1.

Algorithm 1: Localization-Aware Attack on Point Cloud Perturbation

Input: The original clean point cloud X ; number of iterations T ; 3D object detector $f_\theta(\cdot)$.

Output: Final adversarial point cloud X_{adv} .

- 1: Initialize $X_{adv} = X$
 - 2: for $it = 1$ to T do
 - 3: Calculate adversarial loss $J(X, X_{adv}, G)$:
 - 4: $J(X, X_{adv}, G) = -\alpha L_{adv}(f_\theta(X_{adv}), G) + \beta D_C(X, X_{adv})$
 - 5: Update X_{adv} using the gradient of J w.r.t. input
 - 6: end for
 - 7: return X_{adv}
-

Existing adversarial attack methodologies have not sufficiently utilized the spatial features of 3D point clouds to execute attacks against the detector’s localization awareness. Therefore, we propose a novel non-targeted adversarial loss function that generates adversarial samples through gradient-based optimization, which can be formulated as:

$$\min J(X, X_{adv}, G) = -\alpha L_{adv}(f_\theta(X_{adv}), G) + \beta D_C(X, X_{adv}) \quad (1)$$

where L_{adv} is our designed core adversarial loss, D_C is the distance constraint loss. α and β are hyperparameters that balance the two components. The core adversarial loss L_{adv} consists of two parts:

$$L_{adv} = L_{cls}(f_\theta(X_{adv})) + L_{loc}(f_\theta(X_{adv}), G) \quad (2)$$

L_{cls} represents the classification loss of the model (typically Sigmoid Focal Loss). We retain this term to ensure a comprehensive impact on both classification and localization capabilities, and conduct ablation studies in the experimental section to evaluate the effects of its inclusion and exclusion. L_{loc} is our specially designed loss

term tailored for attacking the localization-aware capability of point cloud object detectors. Inspired by DIoU [20] and RDIoU [21], this term comprehensively considers the absolute position of the bounding box, its relative spatial relationship with the ground truth, and angular features in 3D space. It is defined as:

$$L_{loc}(X_{adv}, G) = 1 - IoU - \rho^2/c^2 + \rho_{rot}^2/c_{rot}^2 \quad (3)$$

where the IoU value represents the maximum intersection over union between all predicted boxes and ground truth boxes. ρ^2 and c^2 ensure perturbation generation at the positional level, representing respectively the squared Euclidean distance between the center points of two 3D bounding boxes and the squared diagonal length of the minimum enclosing box that wraps them. ρ_{rot}^2 and c_{rot}^2 further consider the object orientation angle, representing respectively the squared distance between the center points of two bounding boxes under rotation and the squared diagonal length of the minimum rotated enclosing box.

4 Methodology

4.1 Datasets

We adopt the classic KITTI [4] dataset as our dataset. KITTI comprises 14,999 frames of synchronized RGB images and LiDAR point clouds from real-world driving scenarios. KITTI provides precise 3D bounding box annotations for three key dynamic object classes: Car, Pedestrian, and Cyclist, encompassing a total of 80,256 labeled instances. The dataset is split into 7,481 training samples and 7,518 testing samples. Furthermore, KITTI classifies each labeled object into three difficulty levels (Easy, Moderate, Hard) based on factors such as occlusion, truncation, and 2D bounding box height. Our use of the dataset follows the configuration specified in [15].

4.2 Evaluation Metrics

We employ the mAP Ratio as a metric. This is defined as the ratio of the model’s mAP score on adversarial samples to its score on the original clean point clouds. Furthermore, to measure the effectiveness of our localization-aware attack, we introduce the Location Attack Success Rate (ASR). An attack is deemed successful if it causes the IoU between a predicted bounding box and its corresponding ground truth to fall below a predefined threshold K ($K=0.7$ or 0.5). A high-quality adversarial sample must remain as similar as possible to the original point cloud to ensure stealthiness. We adopt the Chamfer Distance (CD) to quantify the magnitude of the perturbation, as Eq. (7).

4.3 Victim Models

We chose the 3D object detection models listed in Table 1 as our victim models. Taking LiDAR point clouds as input, they cover various architectures like one-stage, two-stage, point-based, and voxel-based. This allows us to comprehensively study performance differences between models under attack.

Table 1. Properties of the Victim 3D Object Detectors

Detectors	Representation	Architecture	Params(M)
SECOND [17]	Voxel	1-stage	5.333
Voxel-RCNN [23]	Voxel	2-stage	7.594
PointPillar [1]	Voxel	1-stage	4.841
PV-RCNN [2]	Point-Voxel	2-stage	13.129
PointRCNN [18]	Point	2-stage	4.061

4.4 Parameter Settings

We use an iterative optimization attack with $T=40$ iterations to generate adversarial samples. The hyperparameters α and β are both set to 1. The distance loss weight is 1.0 by default but is reduced to 0.001 for voxel-based detectors. This adjustment prioritizes creating effective perturbations over minimizing perturbation magnitude. The voxel size is set to $[0.05, 0.05, 0.1]$, and the pillar size used by PointPillar is set to $[0.16, 0.16, 4]$. All experiments were conducted on the OpenPCDet platform [22].

5 Results

5.1 Attack Performance

We compare our complete LAA and its ablation variant without classification loss (LAA w/o L_{cls}) against the PGD Perturbation Attack (VPP) proposed in [16].

Table 2 shows the mAP Ratio across five representative 3D object detectors. LAA causes the largest performance drop on all tested models compared to VPP. On PV-RCNN, the detector retains 57.62% of its capability after VPP attack. In contrast, LAA suppresses the mAP ratio to 2.80%. For PointRCNN and PointPillars, LAA renders these detectors nearly non-functional, reducing their mAP ratios to 0.05% and 0.39%, respectively. This is a qualitative improvement over VPP, where PointRCNN still retains 14.96% performance. Even on voxel-based models with higher inherent robustness, such as Voxel-RCNN and SECOND, LAA achieves the lowest mAP ratios of 90.24% and 67.59%, outperforming the baseline in all cases.

Notably, the ablation variant LAA w/o L_{cls} still outperforms VPP in most cases, indicating that the localization-aware mechanism is the key factor behind the attack’s effectiveness.

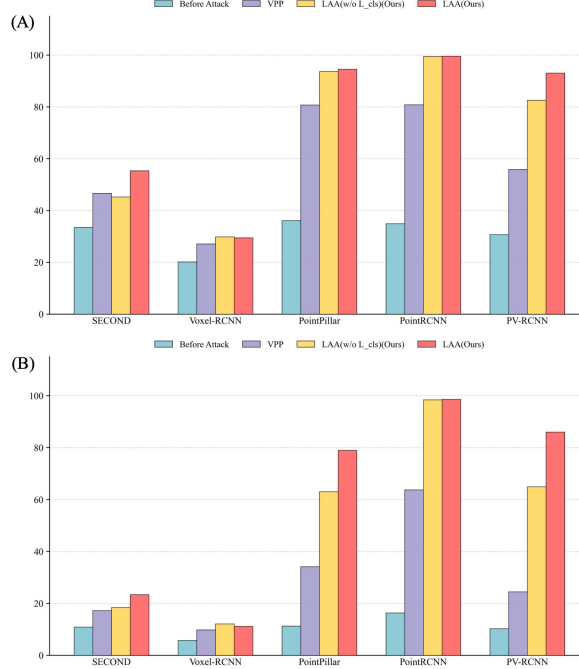
Table 2. mAP Ratio (%) of Victim Models Under Attack

Attack Methods	SECOND	Voxel-RCNN	PointPillar	PointRCNN	PV-RCNN
VPP [16]	68.50	93.96	7.57	14.96	57.62
LAA (ours)	67.59	90.24	0.39	0.05	2.80
LAA w/o L_{cls} (ours)	69.67	91.24	4.61	0.08	10.96

The Location Attack Success Rate (ASR) offers a more direct view of how localization awareness is disrupted. Fig. 2 reports the ASR under two IoU thresholds: the standard threshold ($K=0.7$) and a stricter one ($K=0.5$).

As shown in Fig. 2(A), at $K=0.7$, LAA outperforms VPP on all models. For example, on PV-RCNN, LAA achieves an ASR of 82.57%, compared to 55.86% for VPP. This shows that LAA effectively pushes the IoU of predicted boxes below the acceptable threshold under standard conditions.

The gap widens at $K=0.5$, as shown in Fig. 2(B). At this stricter threshold, a successful attack must force a larger deviation in the bounding box. VPP’s ASR drops noticeably, suggesting it can only induce small perturbations that pass $K=0.7$ but fail at $K=0.5$. In contrast, LAA maintains a high success rate, confirming that it causes substantial deviation between predictions and ground truth. Importantly, the ASR of LAA w/o L_{cls} is nearly the same as the full LAA, which indicates that the localization-aware loss is the primary contributor to attack success, while the classification loss provides only marginal benefit.

**Fig. 2.** ASR results for 3D object detectors under (A) $K=0.7$ and (B) $K=0.5$.

5.2 Evaluation on IoU-Enhanced Detectors

To verify whether relying solely on the overlap metric (IoU) is sufficient to represent 3D spatial semantics, we evaluated our method against SECOND-IoU and PointRCNN-IoU, which incorporate IoU loss during training to enhance geometric feature learning.

As shown in Table 3, although this geometric constraint partially suppresses the effectiveness of the baseline VPP (e.g., VPP’s ASR on PointRCNN-IoU drops to 57.80% at K=0.5), it fails to effectively defend against LAA. Our method successfully breaches this defense mechanism, maintaining a staggering 98.41% ASR on PointRCNN-IoU and consistently outperforming the baseline on SECOND-IoU. This suggests that by explicitly fusing absolute coordinates, relative spatial alignment, and rotation features during optimization, LAA exploits deep geometric "blind spots" that the standard IoU loss cannot cover, thereby effectively dismantling the detector’s localization awareness.

Table 3. mAP Ratio and ASR (%) on IoU-Enhanced Detectors

Model	Method	mAP Ratio	ASR (K=0.7)	ASR (K=0.5)
SECOND-IoU	VPP [16]	42.68	23.11	14.67
	LAA (Ours)	42.91	23.56	15.10
PointRCNN-IoU	VPP [16]	8.24	82.40	57.80
	LAA (Ours)	0.22	99.60	98.41

5.3 Attack Perceptibility

Table 4 presents the quantitative comparison of perceptibility between LAA and VPP using Chamfer Distance (CD) as the metric. The results indicate that LAA maintains a level of imperceptibility comparable to VPP. The VPP method employs strict distance constraints to limit perturbation magnitude; while this "hard constraint" ensures stealthiness, it restricts the optimization of perturbations toward the most destructive direction, sacrificing attack performance. In contrast, LAA achieves significantly superior attack effectiveness while maintaining high stealthiness (similar CD values), demonstrating a more efficient utilization of the perturbation budget.

Table 4. Attack Perceptibility (Chamfer Distance in Meters)

Detector	VPP [16]	LAA (Ours)
SECOND	0.0125	0.0121
Voxel-RCNN	0.0124	0.0127
PointPillar	0.0125	0.0117
PointRCNN	0.0125	0.0118
PV-RCNN	0.0142	0.0119

6 Discussion

6.1 Limitations

While LAA demonstrates superior attack effectiveness across multiple detector architectures, several limitations of our current approach warrant discussion:

- (1) White-box Assumption. Our method requires full access to the detector’s architecture and gradient information for backpropagation-based optimization. Extending LAA to black-box settings, where the attacker has no knowledge of the model internals, remains an open challenge that requires transfer-based or query-based attack strategies.
- (2) Digital Domain Evaluation. The current experiments are conducted entirely in the digital domain by directly perturbing point cloud coordinates. The physical realization of such perturbations (e.g., through adversarial 3D-printed objects or LiDAR spoofing devices) requires additional engineering effort and consideration of environmental factors such as sensor noise and weather conditions.
- (3) Computational Overhead. The iterative gradient-based optimization process incurs non-negligible computational cost, requiring $T = 40$ iterations per sample in our experiments. This overhead may limit the applicability of LAA in real-time attack scenarios where rapid perturbation generation is required.

6.2 Implications for Defense

Our findings reveal critical vulnerabilities in the localization capability of current 3D object detectors, suggesting several directions for building more robust perception systems:

- (1) Geometry-Aware Training. Incorporating rotation-aware and position-aware regularization terms during detector training may enhance robustness against localization-targeted attacks. Explicitly penalizing center offset and angular deviation during training could improve the model’s geometric feature learning.
- (2) Ensemble Methods. Combining detectors with different architectural inductive biases (e.g., point-based and voxel-based representations) may provide complementary robustness. Our results show that voxel-based methods exhibit higher inherent robustness, suggesting that architectural diversity could serve as a defense mechanism.
- (3) Input Validation. Developing statistical methods to detect anomalous perturbation patterns in input point clouds could serve as a preprocessing defense layer. Point cloud denoising or outlier removal techniques may help mitigate the impact of adversarial perturbations.

7 Conclusion and Future Work

Current adversarial attack methods largely overlook the inherent vulnerability of 3D detectors regarding geometric localization. To address this limitation, we propose the

LAA, a point cloud attack method for LiDAR-based 3D object detectors, which is designed to precisely disrupt the model’s localization-aware capability. Experimental results confirm the remarkable effectiveness of this method. We introduce the Location ASR metric to quantify its specific impact, revealing that this targeted localization-aware attack is the dominant factor causing a sharp decline in the model’s comprehensive detection performance (measured by mAP). We not only validate the efficacy of the localization-aware attack strategy but, more importantly, reveal a critical vulnerability in the localization functionality of current 3D object detectors, providing insights for the development of more robust detectors.

For future work, we will focus on establishing more comprehensive evaluation metrics for 3D point cloud spatial features. We will also extend our current attack method to the paradigms of point addition and deletion. Additionally, we will further investigate other potential factors impacting model localization awareness to inform the development of future defense strategies.

Acknowledgment. The authors would like to thank the developers of the OpenPCDet framework for providing the open-source platform used in this research.

References

1. Lang, A.H., Vora, S., Caesar, H., Zhou, L., Yang, J., Beijbom, O.: PointPillars: Fast encoders for object detection from point clouds. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Long Beach, CA, USA, pp. 12697–12705 (2019)
2. Shi, S., Guo, C., Jiang, L., Wang, Z., Shi, J., Wang, X., Li, H.: PV-RCNN: Point-voxel feature set abstraction for 3D object detection. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Seattle, WA, USA, pp. 10529–10538 (2020)
3. Zhou, Y., Tuzel, O.: VoxelNet: End-to-end learning for point cloud based 3D object detection. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Salt Lake City, UT, USA, pp. 4490–4499 (2018)
4. Geiger, A., Lenz, P., Urtasun, R.: Are we ready for autonomous driving? The KITTI vision benchmark suite. In: Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR), Providence, RI, USA, pp. 3354–3361 (2012)
5. Caesar, H., Bankiti, V., Lang, A.H., Vora, S., Liong, V.E., Xu, Q., Krishnan, A., Pan, Y., Baldan, G., Beijbom, O.: nuScenes: A multimodal dataset for autonomous driving. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Seattle, WA, USA, pp. 11621–11631 (2020)
6. Szegedy, C., Zaremba, W., Sutskever, I., Bruna, J., Erhan, D., Goodfellow, I., Fergus, R.: Intriguing properties of neural networks. In: Proc. Int. Conf. Learning Representations (ICLR), Banff, AB, Canada (2014)
7. Goodfellow, I.J., Shlens, J., Szegedy, C.: Explaining and harnessing adversarial examples. In: Proc. Int. Conf. Learning Representations (ICLR), San Diego, CA, USA (2015)
8. Guo, Y., Wang, H., Hu, Q., Liu, H., Liu, L., Bennamoun, M.: Deep learning for 3D point clouds: A survey. *IEEE Trans. Pattern Analysis and Machine Intelligence* 43(12), 4338–4364 (2021)
9. Yan, J., Yin, H.: Adversarial examples in environment perception for automated driving. *arXiv preprint arXiv:2504.08414* (2025)

10. Madry, A., Makelov, A., Schmidt, L., Tsipras, D., Vladu, A.: Towards deep learning models resistant to adversarial attacks. In: Proc. Int. Conf. Learning Representations (ICLR), Vancouver, BC, Canada (2018)
11. Carlini, N., Wagner, D.: Towards evaluating the robustness of neural networks. In: Proc. IEEE Symp. Security and Privacy (SP), San Jose, CA, USA, pp. 39–57 (2017)
12. Xiang, C., Qi, C.R., Li, B.: Generating adversarial point clouds. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Long Beach, CA, USA, pp. 914–922 (2019)
13. Liu, D., Yu, R., Su, H.: Extending adversarial attacks and defenses to deep 3D point cloud classifiers. In: Proc. IEEE Int. Conf. Image Processing (ICIP), Taipei, Taiwan, pp. 2279–2283 (2019)
14. Wang, L., Huang, Y.: A survey of 3D point cloud and deep learning-based approaches for scene understanding in autonomous driving. *IEEE Intell. Transp. Syst. Mag.* 14(6), 135–154 (2022)
15. Zhang, Y., Hou, J., Yuan, Y.: A comprehensive study of the robustness for LiDAR-based 3D object detectors against adversarial attacks. *Int. J. Computer Vision* 132, 1592–1624 (2024)
16. Chen, H., Yan, H., Yang, X., Su, H., Zhao, S., Qian, F.: Efficient adversarial attack strategy against 3D object detection in autonomous driving systems. *IEEE Trans. Intelligent Transportation Systems* 25(11), 16118–16132 (2024)
17. Yan, Y., Mao, Y., Li, B.: SECOND: Sparsely embedded convolutional detection. *Sensors* 18(10), 3337 (2018)
18. Shi, S., Wang, X., Li, H.: PointRCNN: 3D object proposal generation and detection from point cloud. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Long Beach, CA, USA, pp. 770–779 (2019)
19. Tu, C., Li, J., Liu, C., Li, B.: Physically realizable adversarial examples for LiDAR-based 3D object detection. In: Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR), Seattle, WA, USA, pp. 12056–12065 (2020)
20. Zheng, Z., Wang, P., Liu, W., Li, J., Ye, R., Ren, D.: Distance-IOU loss: Faster and better learning for bounding box regression. In: Proc. AAAI Conf. Artif. Intell. 34(07), 12993–13000 (2020)
21. Sheng, H., et al.: Rethinking IoU-based optimization for single-stage 3D object detection. In: *Computer Vision – ECCV 2022*, pp. 544–559. Springer, Cham (2022)
22. OpenPCDet Development Team: OpenPCDet: An open-source toolbox for 3D object detection from point clouds. GitHub repository. <https://github.com/open-mmlab/OpenPCDet> (accessed Sep. 17, 2024)
23. Deng, J., Shi, S., Li, P., Zhou, W., Zhang, Y., Li, H.: Voxel R-CNN: Towards high performance voxel-based 3D object detection. In: Proc. AAAI Conf. Artif. Intell. 35(2), 1201–1209 (2021)