

Descriptor: Context-Aware Collaborative Perception in Autonomous Driving Dataset (ConVeX)

Original

Descriptor: Context-Aware Collaborative Perception in Autonomous Driving Dataset (ConVeX) / Palena, M., Selvaraj, D.C., Chiasserini, C.F., Cerquitelli, T.. - In: IEEE DATA DESCRIPTIONS. - ISSN 2995-4274. - 3:(2026), pp. 531-538. [10.1109/IEEEDATA.2026.3706437]

Availability:

This version is available at: 11583/3010047 since: 2026-06-25T06:24:44Z

Publisher:

IEEE

Published

DOI:10.1109/IEEEDATA.2026.3706437

Terms of use:

This article is made available under terms and conditions as specified in the corresponding bibliographic description in the repository

Publisher copyright

(Article begins on next page)

Received 30 January 2025; revised 13 January 2026, 17 April 2026, and 17 April 2026; accepted 17 April 2026; Date of publication 23 June 2026; date of current version 30 June 2026.

Digital Object Identifier 10.1109/IEEEDATA.2026.3706437

Descriptor: *Context-Aware Collaborative Perception in Autonomous Driving Dataset (ConVeX)*

MARCO PALENA ¹, DINESH CYRIL SELVARAJ ²,
CARLA FABIANA CHIASSERINI ^{1,2,3} (FELLOW, IEEE), AND TANIA CERQUITELLI ⁴

¹CNIT, 43124 Parma, Italy

²CARS@Polito, Politecnico di Torino, 10129 Torino, Italy

³Chalmers University of Technology, SE-41296 Gothenburg, Sweden

⁴Department of Control and Computer Engineering, Politecnico di Torino, 10129 Torino, Italy

CORRESPONDING AUTHOR: Marco Palena (e-mail: marco.palena@polito.it).

This work was supported in part by CARS@Polito, Torino, Italy, the Evolution of Electric Vehicles with V2I, Energy Management, and Smart Mobility (EVOLVE) project (Bando SWich—Programma Regionale Piemonte F.E.S.R. 2021/2027), and PNRR PNC-A.1-N-1—Progetto Living Lab To Move Project—MAAS4ITALY under Grant C15C22007220001, and in part by “WEBFARE” through the Ministero dell’Università e della Ricerca—with the FISA 2022 [D.D. n. 1405 DEL 13/9/2022 Fondo Italiano per le Scienze Applicate (FISA)] program under Grant FISA2022-00908.

ABSTRACT Collaborative perception (CP), using vehicle-to-everything (V2X) communication to share sensor data between connected vehicles and between the latter and the network infrastructure, has emerged as a prominent solution to extend the view of single autonomous vehicles. The effectiveness of this paradigm, however, may be hindered by the presence of adverse weather conditions and changes in lighting, often affecting real-world scenarios. Thus, assessing the robustness of CP to environmental contingencies is still an open issue. Importantly, although some large-scale datasets for CP, comprising realistic and simulated data, are now publicly available, most of them lack diversity in terms of environmental conditions in the autonomous driving scenarios they collect, making it difficult for researchers to assess how such conditions may affect perception performance. We thus introduce Context-Aware Collaborative Perception in Autonomous Driving Dataset (ConVeX), an extensive multi-agent synthetic dataset for CP that reproduces different realistic driving scenarios (urban, rural, and highway), road layouts, and weather and lighting conditions. Remarkably, ConVeX includes multi-modal data [i.e., images from red–green–blue (RGB) cameras, light detection and ranging (LiDAR) points, and global navigation satellite system (GNSS) coordinates] collected by different vehicles, and ground-truth annotations for object detection.

IEEE SOCIETY Vehicular Technology Society (VTS)

DATA TYPE/LOCATION Structured Text Data (JSON); Images (PNG); Point Cloud Data (PCD)

DATA DOI/PID 10.21227/naks-xp87

INDEX TERMS Autonomous vehicles (AVs), collaborative perception (CP), computer vision.

BACKGROUND

Autonomous driving has garnered significant interest from academia and industry due to its potential to deliver safer and more efficient driving experiences. However, most current autonomous driving systems focus on individual vehicles, which limits their environmental perception because of restricted fields of view, thereby continuing to pose safety risks. Collaborative perception (CP)—where autonomous

vehicles (AVs) and road infrastructure devices cooperate through vehicle-to-everything (V2X) communication—has recently emerged as a promising approach to overcoming these limitations. Nevertheless, various factors such as adverse weather conditions and insufficient lighting may still lead to low perception accuracy, exacerbating safety concerns [1], [2]. Weather conditions such as rain, snow, and fog often degrade the color and texture quality of images

while introducing noise that distorts the distribution of light detection and ranging (LiDAR) points. Similarly, low-light environments may cause the loss of details and contrast in images. At the same time, intense direct light from either the sun or artificial light sources may result in overexposed images and induce noise in LiDAR-collected data [3]. These data inaccuracies in environment perception can cause errors or malfunctions in the vehicle's perception system, limiting the applicability of autonomous driving to only specific, highly favorable conditions [4].

Developing solutions to mitigate these issues requires access to comprehensive datasets that encompass a wide range of weather conditions and lighting variations. However, most datasets commonly used to train perception algorithms for autonomous driving lack diversity in these aspects, as they are typically biased toward daytime and clear-weather conditions [3], [5]. As a consequence, recent years have witnessed the emergence of an increasing number of datasets designed to represent a wide range of weather and lighting conditions [5], [6], [7], [8], [9]. However, due to inherent biases and the difficulty of capturing a broad and diverse range of weather and lighting scenarios in real-world settings, these datasets are often synthesized using autonomous driving simulators such as CARLA [10]. Still, these datasets typically collect data from a single sensing agent per frame, making them unsuitable for supporting research on CP. A notable exception is Adver-City [11], a multi-agent, multi-modal synthetic dataset for CP specifically designed to capture a variety of adverse environmental conditions. Furthermore, most studies fail to treat weather conditions and time of day as distinct sources of noise in the perceived data. For instance, variations in lighting are often categorized as separate instances of weather conditions, disregarding the challenges that arise when both factors occur simultaneously—such as a rainy and snowy road at night. Additionally, existing adverse-condition datasets, including synthetic ones, are not free from biases, as specific operating scenarios often occur only in a subset of the overall scenes and, in general, adverse conditions are not systematically varied across the different scenes. Most of these datasets are designed to maximize diversity by augmenting different scenes with adverse weather and lighting conditions. This makes it difficult to assess the extent to which adverse conditions impact the performance of perception algorithms in relation to the scene. In contrast, we focus on a dataset that enables a rigorous assessment of the robustness of CP algorithms under adverse operating scenarios by collecting data from the same scenes across varying conditions. Table I summarizes these observations, highlighting the characteristics of existing datasets for CP and adverse conditions.

To fill the above gaps and overcome the limitations of existing datasets, in this work, we propose Context-Aware Collaborative Perception in Autonomous Driving Dataset (ConVeX), a large-scale, *multi-agent*, synthetic dataset

designed to support V2X CP research with a focus on adverse operating conditions. In contrast to similar works like [11], ConVeX is designed to evaluate collaborative V2X perception robustness under systematically varied adverse weather and lighting, enabling precise ablation studies across a variety of urban, rural, and highway scenarios. Further, we leverage the advanced capabilities of SCANer™ Studio [12], a state-of-the-art industrial autonomous driving simulation platform, to collect images and LiDAR points from multiple vehicular and roadside agents in a variety of virtual scenes, encompassing urban, rural, and highway scenarios.

Unlike existing synthetic datasets, which are typically generated using open-source simulators such as CARLA, our dataset is created using a proprietary, industrial-grade driving simulator widely employed by car manufacturers. Running simulations with this tool requires specialized expertise, significant time, and considerable computing resources. Thus, making this dataset publicly available will greatly benefit the autonomous driving research community. For each scene, we collect data under a wide range of operating conditions by running multiple simulations with various combinations of weather and time of day settings. To the best of our knowledge, ConVeX is the first synthetic, multi-agent dataset for CP specifically designed to address the robustness of CP under adverse operating conditions.

Furthermore, to support downstream perception tasks, such as detection and tracking, we annotate each data sample with 3-D bounding boxes, in addition to labeling each data sample with contextual information such as time of day and weather conditions. Importantly, we also make the raw simulation data publicly available to allow other researchers to extend the dataset with additional annotations.

In summary, ConVeX offers the following key features.

- 1) It provides multi-agent sensory streams that represent each scene as perceived by various vehicles and infrastructure devices.
- 2) It does so in diverse driving scenarios, road layouts, and operating conditions determined jointly by weather and lighting variations, providing data under different sets of conditions for each scene to minimize biases.
- 3) It provides per-vehicle multi-modal data, integrating red–green–blue (RGB) images with LiDAR points and global navigation satellite system (GNSS) coordinates.
- 4) It includes ground-truth annotations for the widely popular task of object detection and can be easily extended to support other tasks as well.
- 5) It provides data obtained via accurate and extensive simulations that would not be otherwise available due to the proprietary nature of the simulator.
- 6) It makes data easily accessible, eliminating the need for expertise in designing relevant scenarios or using the simulator, as well as the expense of substantial computational resources.

TABLE I. Comparison of CP Datasets for Autonomous Driving

Dataset	Source	Comm. Scenario	Goal		Cam	LiDAR	Sensors Depth	GNSS	IMU	Scale			
			Adverse	Robustness						Scenarios	Scenes	Frames	Agents
V2V-Sim [13]	Sim [14]	V2V				✓				> 1	5.5K	51K	1-63
V2X-Sim [15]	Sim [10], [16]	V2X			✓		✓	✓		3	100	10K	2-5
OPV2V [17]	Sim [10], [16], [18]	V2V			✓	✓		✓	✓	8 + 1	73	11K	2-7
DAIR-V2X-C [19]	Real	V2I			✓	✓		✓	✓	1	100	39K	2
V2XSet [20]	Sim [10], [18]	V2X			✓	✓				5	55	11K	2-5
DOLPHINS [21]	Sim [10]	V2X			✓	✓				3	6	42K	3
DAIR-V2X-Seq [22]	Real	V2I			✓	✓		✓		1	95	20K	2
V2V4Real [23]	Real	V2V			✓	✓				1	67	15K	2
RACECAR [24]	Real	none			✓	✓		✓	✓	2	11	50-130K	1-6
WHALES [25]	Sim [10]	V2X			✓	✓				> 5	n.a.	70K	8
TUMTraf-V2X [26]	Real	V2I			✓	✓		✓	✓	1	8	5K	2
RCooper [27]	Real	I2I			✓	✓				2	410	50K	2-4
V2X-Real [28]	Real	V2X			✓	✓		✓	✓	2	68	171K	4
ACDC [5]	Real	none	✓		✓					> 3	n.a.	4K	1
GTA [7]	Sim [29]	none	✓		✓					n.a.	n.a.	300K	1
Adver-City [11]	Sim [10]	V2X			✓			✓	✓	5	110	24K	5
ConVex (ours)	Sim [12]	V2X	✓	✓	✓	✓		✓		12	156	26K	8-9

Note: Column “Comm. scenario” indicates the target communication paradigm, while “Adverse” and “Robustness” indicate, respectively, whether a dataset is targeted at adverse environmental conditions and whether it is designed to assess collaborative perception robustness to adverse conditions.

COLLECTION METHODS AND DESIGN

Collecting autonomous driving data in real-world environments is both time-consuming and resource-intensive. Additionally, field tests are notoriously difficult to control and replicate, often raising safety concerns for operators, especially under adverse conditions. AV simulators provide an appealing alternative for data collection, allowing researchers to create specific scenarios that would be challenging or impossible to reproduce in real-world settings while offering precise control over operating conditions. In this study, we adopt this methodology by using an AV simulator to collect data from a wide range of virtual scenes, each evaluated under diverse operating conditions. Our approach enables the creation of a comprehensive dataset that accurately captures the complexity of real-world scenarios. Below, we explain our rationale for the simulator selection and the methods we used for data collection, processing, and labeling.

Tool Selection and Hardware Setup

Several open-source AV simulation tools exist, with CARLA [10], LG SVL [30], and AirSim [31] standing out as the most feature-rich and widely used. Unlike previous works relying on these academic, open-source tools, we collected data using SCANer™ Studio [12], a professional, comprehensive simulation platform developed by AVSimulation, widely used in the automotive industry for testing and validating vehicle dynamics, advanced driver assistance systems (ADASs), and autonomous driving systems.

Unlike the aforementioned open-source simulators, SCANer™ Studio offers a comprehensive suite of tools and models to create a highly realistic virtual world, encompassing road environments, vehicle dynamics, traffic, sensors, driving behavior, and environmental conditions. Compared to CARLA, which is primarily focused on autonomous driving research and sensor testing in urban environments, SCANer™ Studio offers significantly more accurate models for simulating vehicle dynamics and environmental conditions. This enhanced realism is crucial for generating high-fidelity sensory data, particularly under adverse conditions.

To run our experiments, we used a workstation equipped with an 8-core Intel i7-11700K processor (3.6 GHz), 32 GB of RAM, and an NVIDIA GeForce RTX 3080 Ti GPU with 12 GB of RAM.

Dataset Design

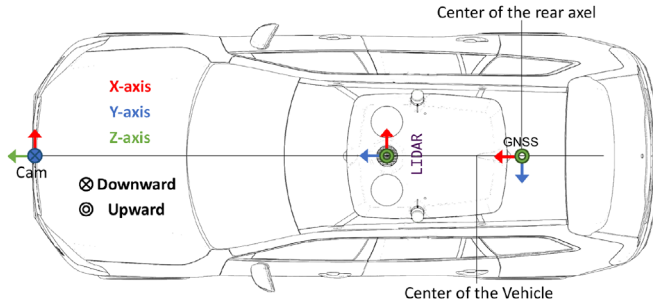
Our data collection strategy is guided by the goal of capturing the same set of virtual scenes across a wide range of operating conditions. Unlike most previous works focusing on adverse-condition datasets, we treat weather (precipitation) and lighting variations (time of day) as independent factors influencing sensor performance, evaluating each scene under every possible combination of these conditions. Our approach enables researchers to assess the impact on perception algorithms of different operating conditions, individually or in combination. For instance, researchers can identify whether performance degradation stems from a specific weather or lighting variation by evaluating their algorithm on data collected from variations of the same base scene.

We set up 13 individual base scenes in SCANer™ Studio, using a mix of pre-existing and custom-made maps. These scenes are designed to capture a wide variety of road layouts and driving conditions in the *urban*, *rural*, and *highway* environments. For each scene, we collect sensory data from multiple agents, including connected autonomous vehicles (CAVs) and roadside units (RSUs). Table II provides a summary of the base scenes used in the ConVex dataset.

As depicted in Fig. 1, each vehicular agent is equipped with a suite of high-resolution RGB camera and LiDAR and GNSS sensors. Each color camera has a resolution of 1600×900 pixels, with 120° horizontal field of view (FOV), 60° vertical FOV, and focal length of 65 mm, recording frames at 20 fps. The simulated LiDAR sensor is mounted vertically on the roof of the vehicle, at a height of 1.8 m from the ground, and replicates the specifications of a Velodyne HDL-64E (64 channels, from -24.8° to $+2^\circ$ of vertical field of view, 20 Hz revolution frequency, 80 m of maximum range and 130K points per frame). Since SCANer™ Studio does not allow for the deployment of fixed sensors within the

TABLE II. Overview of the 13 Base Scenes Used to Create the ConVeX Dataset and Their Distinctive Features: Number of Frames, Annotations, and Agents (Along With Their Relative Average Distance), as Well as Number of Moving and Static Perception Targets and Their Average Number Per Frame

Id	Scenario	Description	No. Frames	No. Annot.	No. Agents		Targets		Avg No. Targets Per Frame	Avg Agent Distance (m)
					CAVs	RSUs	Dyn.	Stat.		
U1_TJ_NTL	Urban	Unregulated T junction	756	4756	6	3	21	47	38	49.73
U4_TJ_TL		Regulated T junction	1716	8994	5	3	25	28	45	74.66
U6_4J_TL		Regulated 4-way junction	2676	18 628	6	3	26	28	42	104.83
U7_4J_TL		Regulated 4-way junction	3468	14 960	6	3	18	4	21	51.77
U8_RA_S		Roundabout small	1284	6545	6	3	19	6	25	39.32
U9_RA_L		Roundabout large	2316	6050	5	3	17	24	39	86.66
U10_SS_X		Straight segment, crosswalk	1776	8164	7	2	17	2	18	122.55
U11_SS_C		Straight segment, construction site	1536	4973	7	2	18	3	21	104.72
R1_2LR	Rural	2-lanes road	3180	19 216	6	3	21	7	35	233.38
R3_NWR		Narrow and winding road	3408	5444	6	3	17	1	12	207.42
H4_TB	Highway	Toll booths	612	3384	7	2	19	32	40	268.62
H5_EE		Highway entering	1536	3057	7	2	15	21	20	160.54
H6_EE		Highway exiting	1824	5999	6	3	18	22	30	222.34

**FIG. 1.** Layout of sensors for vehicular agents.

road infrastructure, we implemented RSUs as a workaround by attaching sensors to stationary pedestrian actors, which will be invisible in the simulations. The sensors of the infrastructural agents are positioned 15 m above the ground, looking diagonally downward at 35° and oriented toward key points of interest in the road infrastructure, such as intersections and toll booths. Each RSU is equipped with a color camera and a LiDAR sensor, both sharing the same global coordinates, and their specifications are identical to those of the vehicular agents. The sensor layout is identical across all vehicles and RSUs.

Regarding weather conditions, we consider four scenarios: *clear*, *cloudy*, *rainy*, and *snowy*. Regarding the time of day, we differentiate between twilight, daylight, and nighttime settings, corresponding to 6:00 a.m., 12:00 p.m., and 8:00 p.m., respectively. Fig. 2 shows a frame taken from a sample scene under all possible combinations of weather and lighting conditions.

Collection Methodology

For each base scene, we ran individual simulations collecting sensory data under all combinations of weather and time of day, thus obtaining a total of 156 scenes. After synchronizing each sensor in SCANer™ Studio, we ran each simulation, collecting data at a frequency of 20 Hz, resulting in each

scene being composed of number of frames varying between 756 and 3468. For each frame, we extract video streams from the RGB camera of the agents as well as LiDAR point clouds. In addition, we collect GNSS data from all vehicles appearing in the scene. We also instructed SCANer™ Studio to generate a list of ground-truth bounding-box annotations for each potential detection target in every frame. To support a variety of downstream detection tasks, we included dynamic and static objects in our target list. Both types of objects are categorized into three main classes, namely *vehicles*, *vulnerable vehicles*, and *pedestrians* for dynamic objects, and *traffic sign*, *traffic signal*, and *barrier* for static objects. Each main class is further divided into several subclasses. Fig. 3 illustrates the target classes breakdown and their occurrences in the ConVeX dataset.

Bounding boxes are derived from ground-truth annotations generated directly by the simulation environment. First, the 3-D coordinates of the eight corner points of each object's bounding box, expressed in the global reference frame of the simulation, are extracted from SCANer™ Studio. They are then post-processed to obtain a compact representation consisting of the 3-D position of the box center (x, y, z) , its dimensions (w, l, h) , and its 3-D orientation (q_w, q_x, q_y, q_z) . We then filter these annotations by removing the bounding boxes that contain no LiDAR points. Finally, we filter out those without any suitable annotations, yielding 26 088 annotated frames, each comprising the sample image and point cloud data for each agent.

VALIDATION AND QUALITY

We now present scene statistics and the generated annotations to emphasize ConVeX data diversity and quality.

Scenes Statistics

Our scenes feature vehicles traveling at different speeds, from stationary up to a maximum of 136 km/h. Fig. 4(a) illustrates the vehicle speed distribution in the urban, rural, and highway scenarios, where the speed values have been

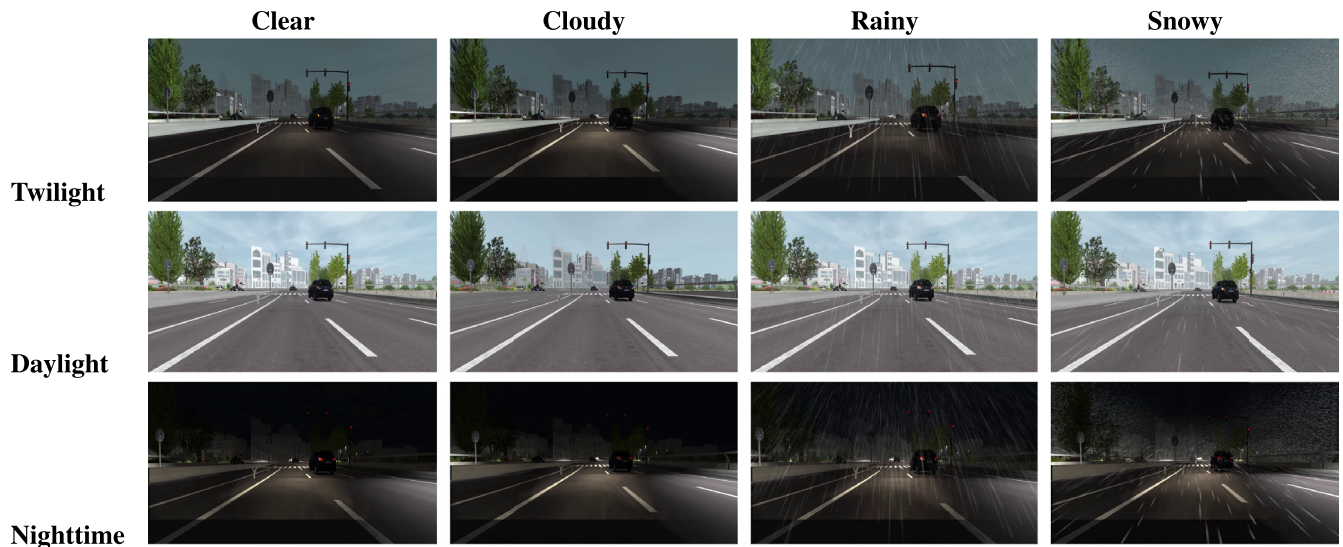


FIG. 2. Sample scene depicted under different operating conditions.

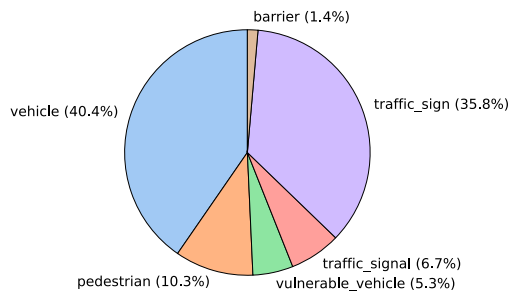


FIG. 3. Distribution of main target classes in the ConVeX dataset.

computed by considering only vehicles that are within 70 m from at least another agent and by averaging, for each of such vehicles, its speed over all scenes and frames. We observe that, consistently with real-world behavior, urban scenes, primarily including vehicles approaching intersections, are characterized by lower speeds, mostly between 20 and 30 km/h. In rural and highway scenes, instead, vehicles tend to travel faster, with speeds mainly between 30 and 40 km/h in the rural scenario and higher than 90 km/h on highways.

The dataset includes scenes with varying vehicle density. Fig. 4(b) illustrates the distribution of the minimum Euclidean distance between each pair of vehicles within a 70-meter range from at least one agent. Again, distance values are averaged across all scenes and frames. Urban scenes tend to be more crowded, with most vehicles located within 20 m of another vehicle. Highway scenes, while less dense than urban ones, still show quite high values of vehicular density, with most vehicles falling within a 20 to 30-meter range from at least one other vehicle. This is a direct consequence of the specific considered highway scenarios (i.e., vehicles approaching a toll booth or entering/exiting the highway via ramps. In contrast, rural scenes exhibit lower vehicle density,

with some vehicles separated by as much as 90 m from the nearest vehicle.

Fig. 4(c) shows the distribution of minimum longitudinal distances between CAV agents and other vehicles traveling on the same road and lane, providing a more realistic measure of inter-vehicle spacing in autonomous driving scenarios. Compared to Euclidean distances, all scenarios shift toward larger separations: urban traffic peaks around 20–30 m, while rural and highway traffic are more spread out, with most vehicles maintaining 30–40 m headways, reflecting free-flow conditions and lane-aligned spacing.

Annotation Statistics

Fig. 5 presents the occurrence of the number of annotations per frame, with most frames containing three to six bounding boxes. Fig. 6, instead, highlights the distribution of the size of the bounding boxes across different classes of target objects. The wide range of vehicle sizes underscores the broad spectrum of real-world vehicles that is represented in our database.

Benchmark

To demonstrate how the ConVeX dataset can be used to assess the robustness of CP algorithms under adverse conditions, we provide a preliminary benchmark using CoBEVT [32], a multi-agent, multi-camera birds eye view (BEV) semantic segmentation model. We first split the ConVeX dataset into train/validation/test sets using a 60/20/20 ratio, yielding 15 652/5218/5218 frames per set. The split is stratified across scenes to ensure a representative mix of urban, rural, and highway scenarios.

We evaluate two variants of CoBEVT on the test set. The first (CoBEVT-V2X-Sim) is the CoBEVT model trained on the V2X-Sim dataset (consisting mostly of clear weather and daytime scenes) as in [33]. The second (CoBEVT-ConVeX)

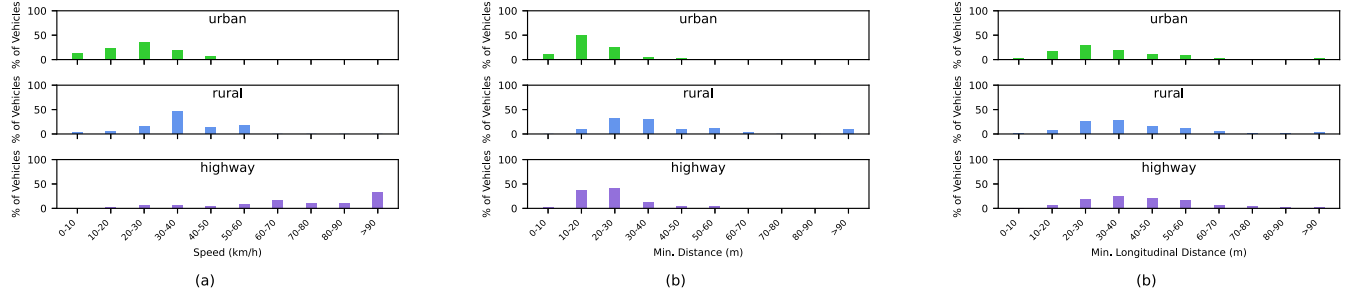


FIG. 4. (a) Distribution of average vehicle speed, (b) minimum inter-vehicle distance, and (c) minimum longitudinal distance between each CAV agent and the other vehicles on the same lane for each scenario, computed considering only those vehicles that are within a 70-meter range from at least one agent.

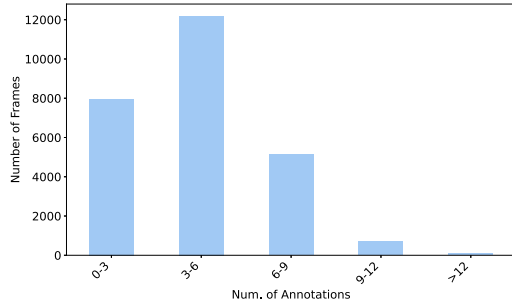


FIG. 5. Occurrence of the number of annotations per frame.

is obtained by fine-tuning the first model on the ConVeX training set for 30 additional epochs with learning rate of 0.00001. All results are reported in terms of AP@50, i.e., average precision computed at an intersection-over-union (IoU) threshold of 0.5. Table III presents the results, broken down by operating conditions. Consistent with expectations, results for both models demonstrate the negative impact of adverse operating conditions on the perception task, with clear weather and daytime scenarios achieving the highest scores, while snowy nighttime conditions emerge as the most challenging. Comparing the two models, CoBEVT-V2X-Sim achieves AP@50 scores ranging from 28 to 44, with a maximum gap of 15.8. Fine-tuning with ConVeX leads to substantial performance improvements across all scenarios, ranging from +13.77% to +32.3%, with a maximum gap of 12.7. These results indicate that fine-tuning on the ConVeX dataset effectively enhances the robustness of CoBEVT under adverse conditions. Notably, the largest gains are observed in rainy and snowy scenarios, which align with the fact that the V2X-Sim dataset consists predominantly of clear weather and daytime scenes.

RECORDS AND STORAGE

The collected data are stored in a modified version of the *nuScenes* [34] storage format, adapted to support multiple agents and varying operating conditions per frame. As a result, the dataset can be parsed using the *nuScenes devkit* [35]. Our storage format differs from the original *nuScenes* format in the following ways.

- 1) We do not store map data (e.g., images and semantic maps).
- 2) We do not record the visibility level of instances (e.g., the fraction of a detection target appearing in an image).
- 3) We do record the base scene and the operating conditions associated with each scene, by encoding both pieces of information in the description field of the scene, using the following syntax: `BASE_SCENE; WEATHER; TIME_OF_DAY; DESCR`, where `BASE_SCENE` is the base scene identifier as reported in Table II, `WEATHER` and `TIME_OF_DAY` are labels describing the scene's operating conditions, as reported in Fig. 2, and `DESCR` contains any additional description of the scene.

All annotations and metadata—including sensor parameters, maps, ego poses, and operating conditions—are stored in JavaScript object notation (JSON) files and organized within a relational database, as detailed in the *nuScenes* documentation [35]. Sensor data, annotations, and metadata are structured in files and folders as outlined in Table IV.

INSIGHTS AND NOTES

We remark that, by providing both the raw data and the source code, we enable other researchers to expand the ConVeX dataset with additional annotations, facilitating the support for further downstream tasks such as semantic segmentation.

SOURCE CODE AND SCRIPTS

The source code used to generate the dataset includes a collection of Python scripts for the sampling, processing, and annotation of raw data and for the generation of the final samples in the modified *nuScene* format described in “Records and Storage” section. Scripts take as input the raw data we generated for each scene. Such data are organized in a folder hierarchy, with folders labeled by the base scene identifier at the top level (level 0), the weather conditions at level 1, and the time of day at level 2. Each scene in level-2 folder includes two subfolders: `data` and `videos`. The former contains GNSS data (`gpsData.csv`), LiDAR data (`lidarData.csv`), and ego pose data (`vehData.csv`).

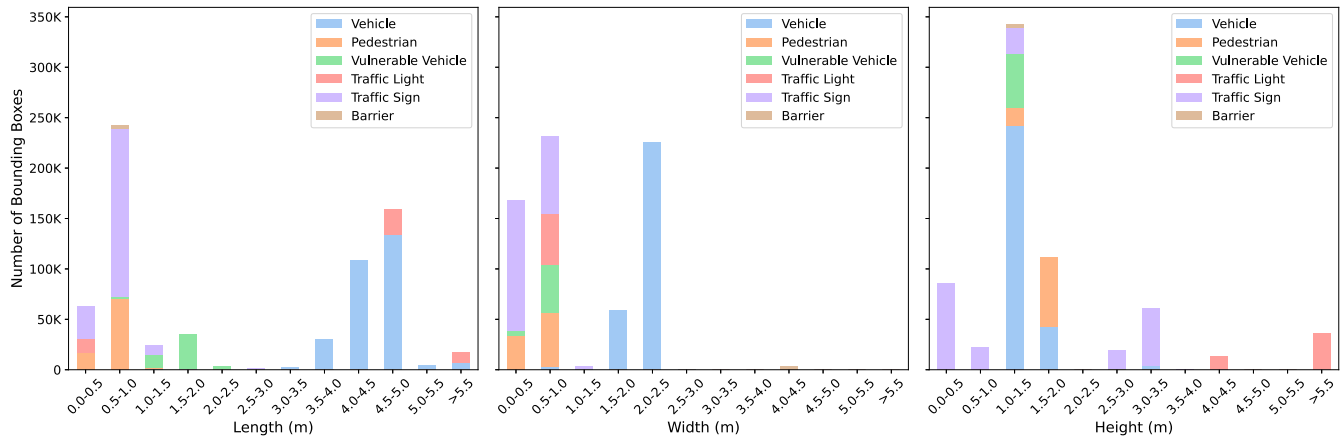


FIG. 6. Distribution of the length, width, and height of bounding box annotations across different classes of dynamic and static objects.

TABLE III. Robustness Assessment of the CoBEVT CP Algorithm Under Adverse Conditions, Using the ConVeX Dataset

Time	CoBEVT-V2X-Sim				CoBEVT-ConVeX			
	Clear	Cloudy	Rainy	Snowy	Clear	Cloudy	Rainy	Snowy
Twilight	40.32	39.61	30.08	32.24	46.38 (+15.0%)	46.55 (+17.5%)	37.85 (+25.8%)	39.27 (+21.8%)
Daytime	44.00	41.08	34.18	34.46	50.03 (+13.7%)	47.36 (+15.3%)	41.72 (+22.1%)	42.58 (+23.6%)
Nighttime	38.15	35.37	31.02	28.21	44.83 (+17.5%)	42.38 (+19.8%)	40.06 (+29.2%)	37.32 (+32.3%)

Note: CoBEVT-V2X-Sim: CoBEVT model trained on the V2X-Sim dataset. CoBEVT-ConVeX: obtained by fine-tuning CoBEVT-V2X-Sim on ConVeX. All results are reported in terms of AP@50, with percentage increase relative to CoBEVT-V2X-Sim between parentheses.

TABLE IV. Folder Structure of the ConVeX Dataset

File or Folder	Description
sampled/	Folder for sensor data at key-frames.
sampled/frames	Folder for sampled video frames (PNG files).
sampled/gnss	Folder for sampled GNSS data (NPY files).
sampled/points	Folder for sampled point clouds data (PCD files).
v1.0-ConVeX/	Folder for dataset annotations and metadata (JSON files).
v1.0-ConVeX/scene.json	Sequences of consecutive frames, from X to Y seconds long, extracted from a log.
v1.0-ConVeX/sample.json	Annotated snapshots of a scene at a particular timestamp (frames).
v1.0-ConVeX/sample_data.json	Data collected from a particular sensor at a given frame.
v1.0-ConVeX/sample_annotation.json	Bounding boxes defining the position of an object seen in a sample.
v1.0-ConVeX/instance.json	Instances of detection target classes observed in samples.
v1.0-ConVeX/category.json	Taxonomy of detection target classes.
v1.0-ConVeX/attribute.json	Properties of an instance that can change while the category remains the same.
v1.0-ConVeX/sensor.json	Sensor types.
v1.0-ConVeX/calibrated_sensor.json	Definition of particular sensors (camera/LiDAR) as calibrated on a particular vehicle.
v1.0-ConVeX/ego_pose.json	Ego vehicle poses at a particular timestamp.
v1.0-ConVeX/log.json	Logs/recordings from which the data was extracted.

for every agent in the scene, all in comma-separated values (CSV) format. The latter contains the video feeds (in audio video interleave (AVI) format) for each agent.

Both the scripts and the raw data used to generate the ConVeX dataset are publicly available at [36], which also provides additional documentation on how to navigate the data and run the scripts.

ACKNOWLEDGMENT

This article reflects only the authors' views and opinions, and the Ministry cannot be considered responsible for them.

M.P. contributed to the scene and dataset design, data annotation, and paper writing. D.C.S. designed the experiments

to obtain the dataset. C.F.C. and T.C. supervised the work and contributed to the paper writing.

The authors declare that they have no conflicts of interest.

REFERENCES

- [1] Y. Zhang, A. Carballo, H. Yang, and K. Takeda, "Perception and sensing for autonomous vehicles under adverse weather conditions: A survey," *ISPRS J. Photogramm. Remote Sens.*, vol. 196, pp. 146–177, Feb. 2023.
- [2] J. Wang, Z. Wu, Y. Liang, J. Tang, and H. Chen, "Perception methods for adverse weather based on vehicle infrastructure cooperation system: A review," *Sensors*, vol. 24, no. 2, Jan. 2024, Art. no. 374.
- [3] A. Carballo et al., "LIBRE: The multiple 3D LiDAR dataset," in *Proc. IEEE Intell. Veh. Symp. (IV)*, Las Vegas, NV, USA: IEEE, Oct. 2020, pp. 1094–1101.
- [4] M. J. Mirza et al., "Robustness of object detectors in degrading weather conditions," in *Proc. IEEE Int. Intell. Transport. Syst. Conf. (ITSC)*, 2021, pp. 2719–2724.

- [5] C. Sakaridis, D. Dai, and L. Van Gool, "ACDC: The adverse conditions dataset with correspondences for semantic driving scene understanding," in *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*, Oct. 2021, pp. 10765–10775.
- [6] W. Madder, G. Pascoe, C. Linegar, and P. Newman, "1 Year, 1000km: The Oxford RobotCar dataset," *Int. J. Robot. Res. (IJRR)*, vol. 36, no. 1, pp. 3–15, 2017. [Online]. Available: <http://dx.doi.org/10.1177/0278364916679498>
- [7] D. Liu, Y. Cui, Z. Cao, and Y. Chen, "A large-scale simulation dataset: Boost the detection accuracy for special weather conditions," in *Proc. Int. Joint Conf. Neural Netw. (IJCNN)*, Glasgow, U.K.: IEEE, Jul. 2020, pp. 1–8.
- [8] L. H. Pham, D. N.-N. Tran, and J. W. Jeon, "Low-light image enhancement for autonomous driving systems using DriveRetinex-Net," in *Proc. IEEE Int. Conf. Consumer Electron. – Asia (ICCE-Asia)*, Seoul, South Korea: IEEE, Nov. 2020, pp. 1–5.
- [9] S. Sural, N. Sahu, and R. R. Rajkumar, "ContextualFusion: Context-based multi-sensor fusion for 3D object detection in adverse operating conditions," in *Proc. IEEE Intell. Veh. Symp. (IV)*, Jeju Island, Republic of Korea: IEEE, Jun. 2024, pp. 1534–1541.
- [10] A. Dosovitskiy, G. Ros, F. Codevilla, A. Lopez, and V. Koltun, "CARLA: An open urban driving simulator," in *Proc. 1st Annu. Conf. Robot Learn.*, 2017, pp. 1–16.
- [11] M. Karvat and S. Givigi, "Adver-City: Open-source multi-modal dataset for collaborative perception under adverse weather conditions," Mar. 2025, pp. 740–746.
- [12] A. Simulations, "Scanner studio," 2024. Available: <https://www.avsimulation.com>
- [13] B. Yang, W. Luo, and R. Urtasun, "PIXOR: Real-time 3D object detection from point clouds," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.*, Jun. 2018, pp. 7652–7660.
- [14] S. Manivasagam et al., "LiDARsim: Realistic LiDAR simulation by leveraging the real world," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit., CVPR*, Seattle, WA, USA: Piscataway, NJ, USA: IEEE, Jun. 13–19, 2020, pp. 11164–11173.
- [15] Y. Li et al., "V2X-Sim: Multi-agent collaborative perception dataset and benchmark for autonomous driving," *IEEE Robotics and Automation Letters*, vol. 7, no. 4, Jul. 2022.
- [16] P. A. Lopez et al., "Microscopic traffic simulation using sumo," in *Proc. 21st IEEE Int. Conf. Intell. Transport. Syst.*, Piscataway, NJ, USA: IEEE, 2018, pp. 2575–2582. [Online]. Available: <https://elib.dlr.de/124092/>
- [17] R. Xu, H. Xiang, X. Xia, X. Han, J. Li, and J. Ma, "OPV2V: An open benchmark dataset and fusion pipeline for perception with vehicle-to-vehicle communication," in *Proc. Int. Conf. Robot. Automat. (ICRA)*, May 2022, pp. 2583–2589.
- [18] R. Xu, Y. Guo, X. Han, X. Xia, H. Xiang, and J. Ma, "OpenCDA: An open cooperative driving automation framework integrated with co-simulation," in *Proc. IEEE Int. Intell. Transport. Syst. Conf. (ITSC)*, Piscataway, NJ, USA: IEEE, 2021, pp. 1155–1162.
- [19] H. Yu et al., "DAIR-V2X: A large-scale dataset for vehicle-infrastructure cooperative 3D object detection," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jun. 2022, pp. 21329–21338.
- [20] R. Xu, H. Xiang, Z. Tu, X. Xia, M.-H. Yang, and J. Ma, "V2X-ViT: Vehicle-to-everything cooperative perception with vision transformer," in *Proc. 17th Eur. Conf. Comput. Vis., ECCV*, Tel Aviv, Israel, Part XXXIX, Berlin, Heidelberg: Springer-Verlag, Oct. 23–27, 2022, pp. 107–124.
- [21] R. Mao, J. Guo, Y. Jia, Y. Sun, S. Zhou, and Z. Niu, "DOLPHINS: Dataset for collaborative perception enabled harmonious and interconnected self-driving," in *Proc. Asian Conf. Comput. Vis., ACCV*, L. Wang, J. Gall, T.-J. Chin, I. Sato, and R. Chellappa, Eds. Cham, Switzerland: Springer Nature, 2023, pp. 495–511.
- [22] H. Yu et al., "V2X-Seq: A large-scale sequential dataset for vehicle-infrastructure cooperative perception and forecasting," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jun. 2023, pp. 5486–5495.
- [23] R. Xu et al., "V2V4Real: A real-world large-scale dataset for vehicle-to-vehicle cooperative perception," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jun. 2023, pp. 13712–13722.
- [24] A. Kulkarni et al., "RACECAR—The dataset for high-speed autonomous racing," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst. (IROS)*, Oct. 2023, pp. 11458–11463.
- [25] Y. Wang, S. Chen, Z. Song, and S. Zhou, "WHALES: A multi-agent scheduling dataset for enhanced cooperation in autonomous driving," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst. (IROS)*, Aug. 2025, pp. 20487–20493.
- [26] W. Zimmer, G. A. Wardana, S. Sritharan, X. Zhou, R. Song, and A. C. Knoll, "TUMTraF V2X cooperative perception dataset," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Mar. 2024, pp. 22668–22677.
- [27] R. Hao et al., "RCooper: A real-world large-scale dataset for roadside cooperative perception," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Seattle, WA, USA: IEEE, Jun. 2024, pp. 22347–22357.
- [28] H. Xiang et al., "V2X-Real: A large-scale dataset for vehicle-to-everything cooperative perception," in *Proc. Eur. Conf. Comput. Vis., ECCV*, A. Leonardis, E. Ricci, S. Roth, O. Russakovsky, T. Sattler, and G. Varol, Eds., Cham, Switzerland: Springer Nature, 2025, pp. 455–470.
- [29] S. R. Richter, V. Vineet, S. Roth, and V. Koltun, "Playing for data: Ground truth from computer games," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, LNCS, B. Leibe, J. Matas, N. Sebe, and M. Welling, Eds., vol. 9906. Cham, Switzerland: Springer International Publishing, 2016, pp. 102–118.
- [30] G. Rong et al., "LGSVL simulator: A high fidelity simulator for autonomous driving," 2020, *arXiv:2005.03778*.
- [31] S. Shah, D. Dey, C. Lovett, and A. Kapoor, "AirSim: High-fidelity visual and physical simulation for autonomous vehicles," in *Field and Service Robotics*, M. Hutter and R. Siegwart, Eds. Cham: Springer International Publishing, 2018, pp. 621–635.
- [32] R. Xu, Z. Tu, H. Xiang, W. Shao, B. Zhou, and J. Ma, "CoBEVT: Cooperative bird's eye view semantic segmentation with sparse transformers," in *Proc. Conf. Robot Learn. (CoRL)*, 2022, pp. 989–1000.
- [33] Y. Lu, Y. Hu, Y. Zhong, D. Wang, S. Chen, and Y. Wang, "An extensible framework for open heterogeneous collaborative perception," in *Proc. 12th Int. Conf. Learning Represent.*, 2024, pp. 44576–44593.
- [34] H. Caesar et al., "nusenes: A multimodal dataset for autonomous driving," 2019, *arXiv:1903.11027*.
- [35] Motional Inc., "nusenes™ devkit." Accessed: Apr. 17, 2026. [Online]. Available: <https://github.com/nutonomy/nusenes-devkit>
- [36] M. Palena, D. C. Selvaraj, C. F. Chiasserini, and T. Cerquitelli, "Convex github repository." Accessed: Apr. 17, 2026. [Online]. Available: <https://github.com/marcopalena/convex/>