



POLITECNICO
MILANO 1863

DEPARTMENT OF ELECTRONICS
INFORMATION AND BIOENGINEERING

From Ego to Eco -System: Perception & Sensor Fusion for Autonomous and Connected Mobility

11.11.2025 | Simone Mentasti

Me, myself, and I

Simone Mentasti, PhD
Researcher
Artificial Intelligence and
Robotics Lab
Dept. Of Electronics,
Information & Bioengineering
Politecnico di Milano
simone.mentasti@polimi.it



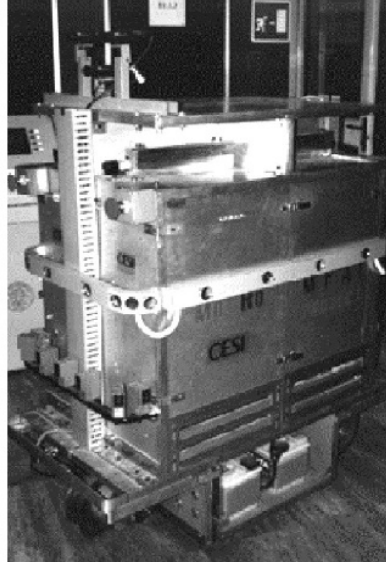
My research interests

- Mobile Robotics
- Autonomous Systems
- Computer Vision & Perception
- Deep Learning

Courses I tech

- Computer Science (BS)
- Robotics (BS + MS)

AIRLab @PoliMi



5 Full Professors
5 Ass. Professors
7 Researchers
44 PhD Students
18 Research Ass.



Contents

1. Autonomous Vehicles Perception

Line detection

Line mapping

Road Users detection and tracking

2. Cooperative Connected Vehicles Perception

Infrastructure -based detection and tracking

Low power-real time optimization

3. Other ongoing projects

Autonomous Vehicles Perception

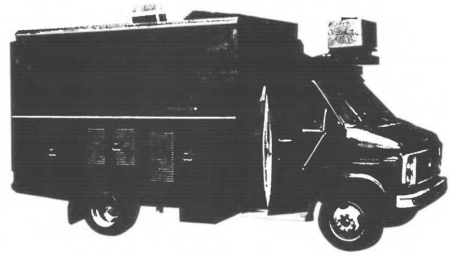
01

Autonomous Vehicles Perception

(Mostly line detection)

01

Autonomous vehicles timeline



ALVINN



Sandstorm

Nissan AV prototype



Wyamo first public taxi



Darpa ALV



Sandstorm



Stanford Racing Junior

Wyamo/google av



Indy autonomous



1985

1988

1990

2001

2007

2013

2015

2020

2021

Sensor setup

Lots of sensors



Sensor setup

Lots of sensors



Sensor setup

Minimal setup



Sensor setup

Minimal setup

Ego vehicle state estimation

GPS is not enough:

- Tunnels
- Multipath
- Urban canyon

Combine other data

Leverage maps



Obstacle detection and tacking

Multiple sensors

- Different data
- Different accuracy
- Different reliability

Retrieve uniform representation of vehicle surroundings

Am I on the road?

Road Line detection

512x512 RGB images from forward facing camera



Datasets

BDD100k



CULane



Pan, X., Shi, J., Luo, P., Wang, X. and Tang, X., 2018, April. Spatial as deep: Spatial pyramid pooling for traffic scene understanding. In *Proceedings of the AAAI conference on artificial intelligence* (Vol. 32, No. 1).

Yu, F., Chen, H., Wang, X., Xian, W., Chen, Y., Liu, F., Madhavan, V. and Darrell, T., 2020. Bdd100k: A diverse driving dataset for heterogeneous multitask learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition* (pp. 2636-2645).

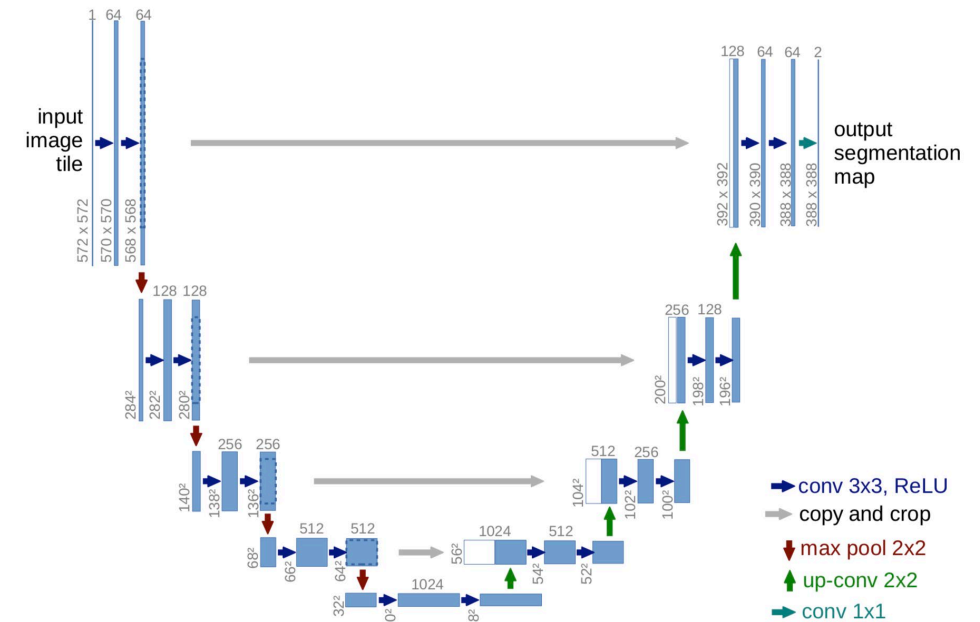
Am I on the road?

Road Line detection

512x512 RGB images from forward facing camera



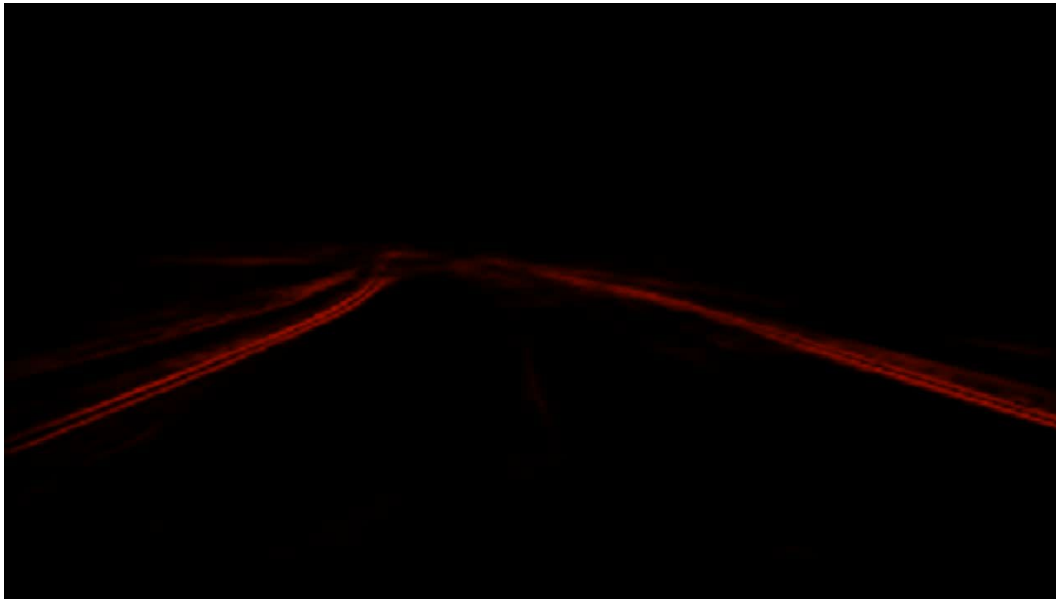
U-Net style segmentation model
Low-power optimized
100Hz on Jetson Xavier



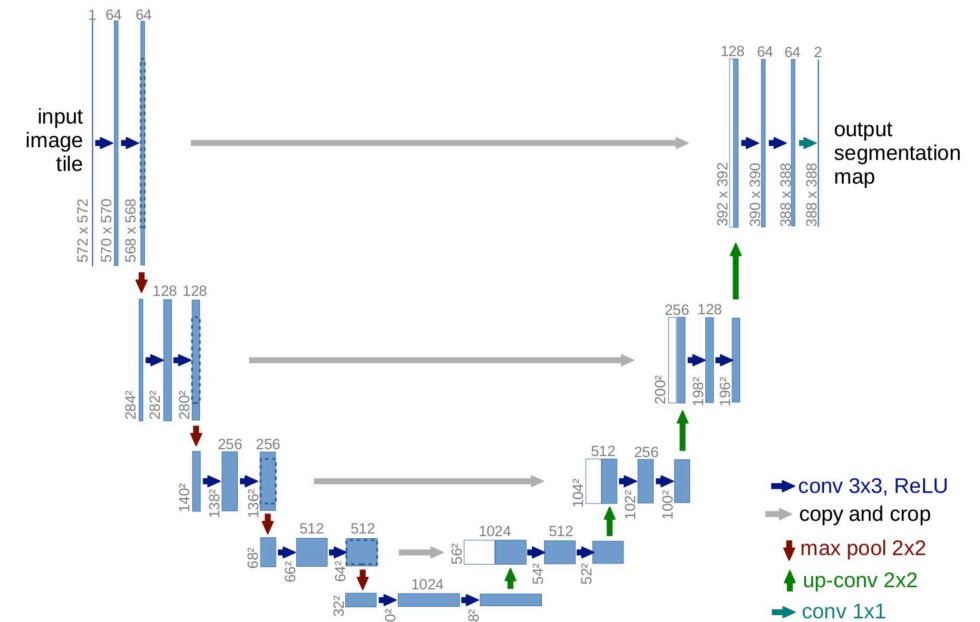
Road lines detection

Deep Learning Based

512x512 RGB images from forward facing camera



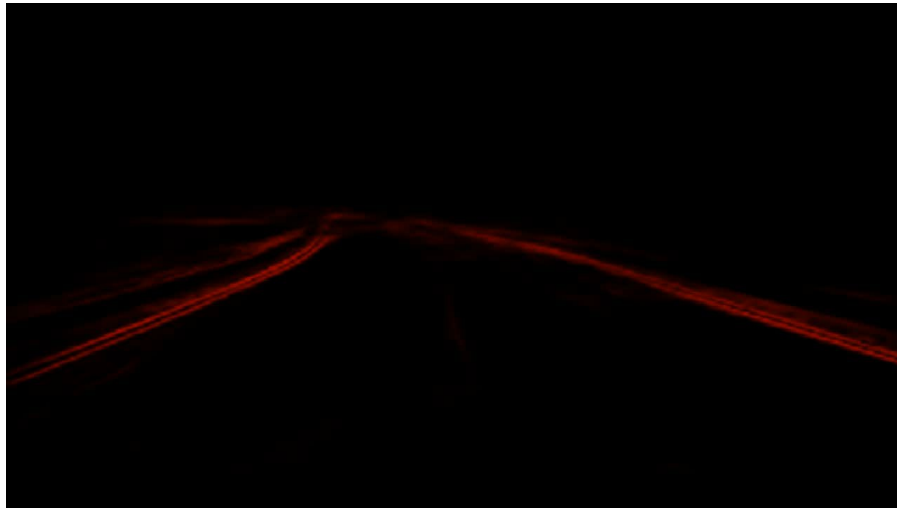
U-Net style segmentation model
Low-power optimized
100Hz on Jetson Xavier



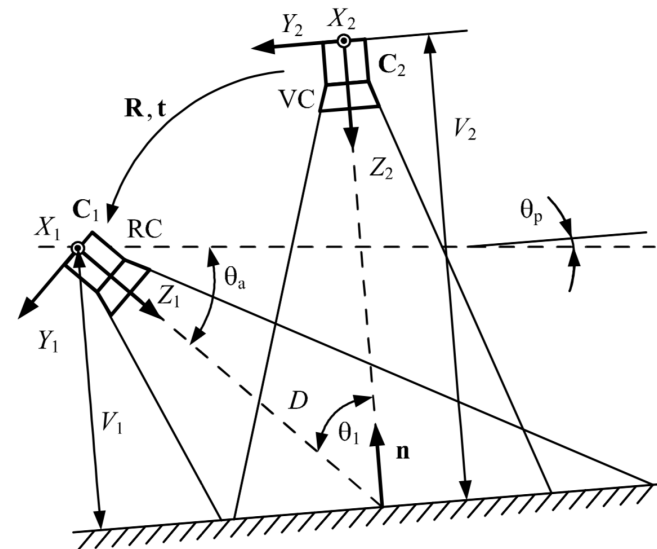
Road lines detection

BEV Reprojection

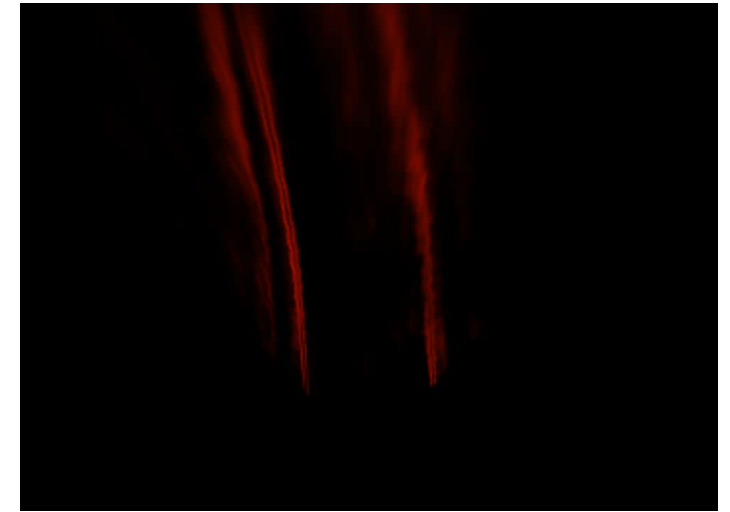
Segmented image



Calibrated Camera



Bird's eye view



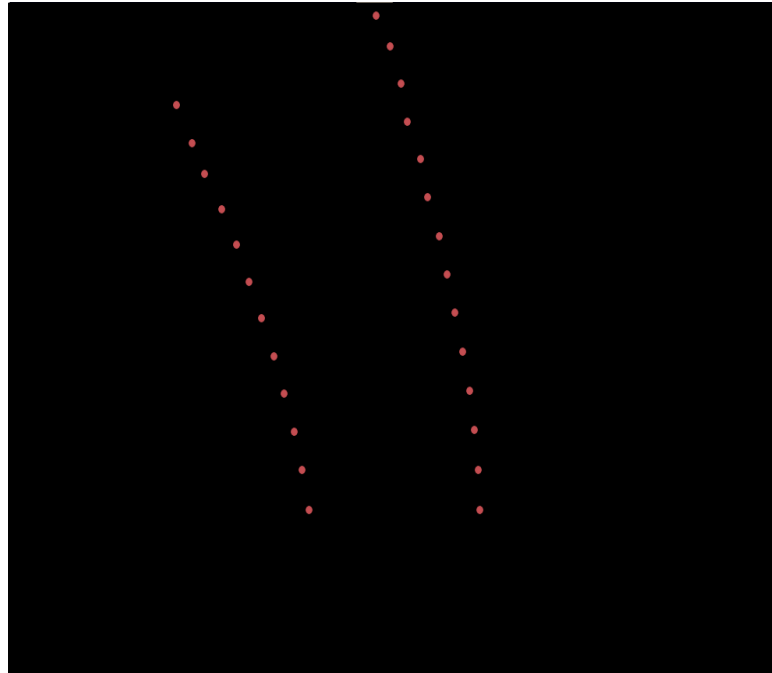
Road lines detection

Fitting

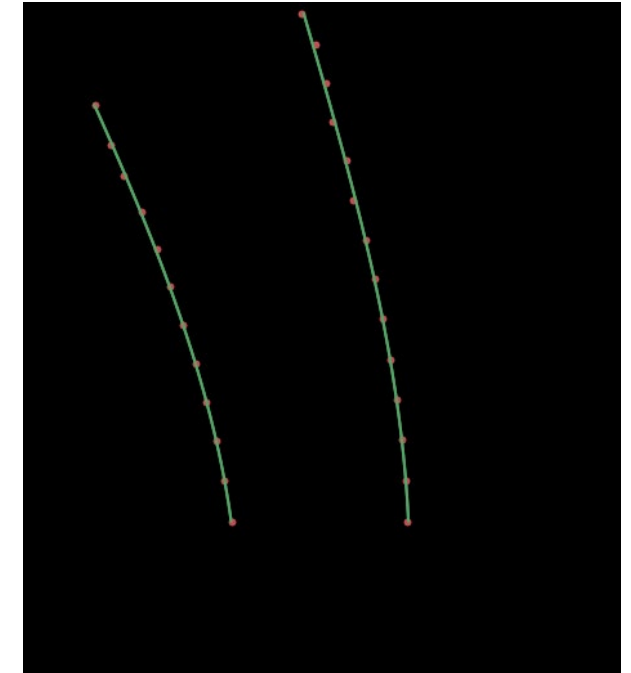
Thresholding (probability)



Point sampling



Line fitting



Road lines position

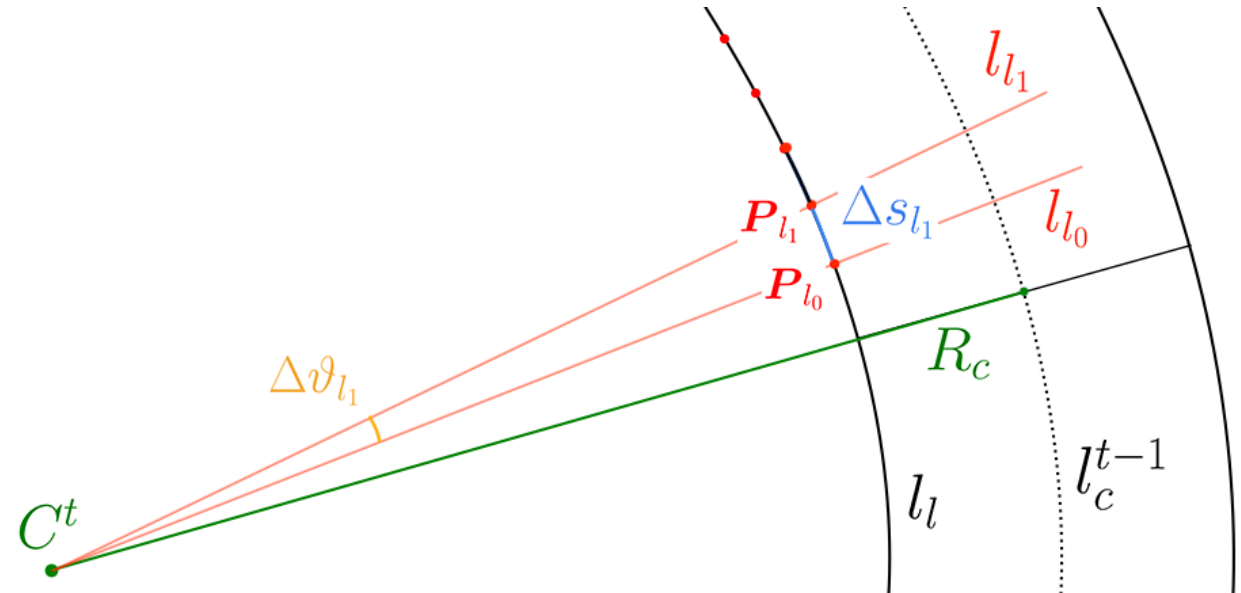
Locally constant curvature

The two lateral lines share the same center of curvature

Center of curvature varies smoothly

Under these assumptions it is possible to project lateral line points to the centreline

Centreline is fitted with a polynomial in (s, ϑ)

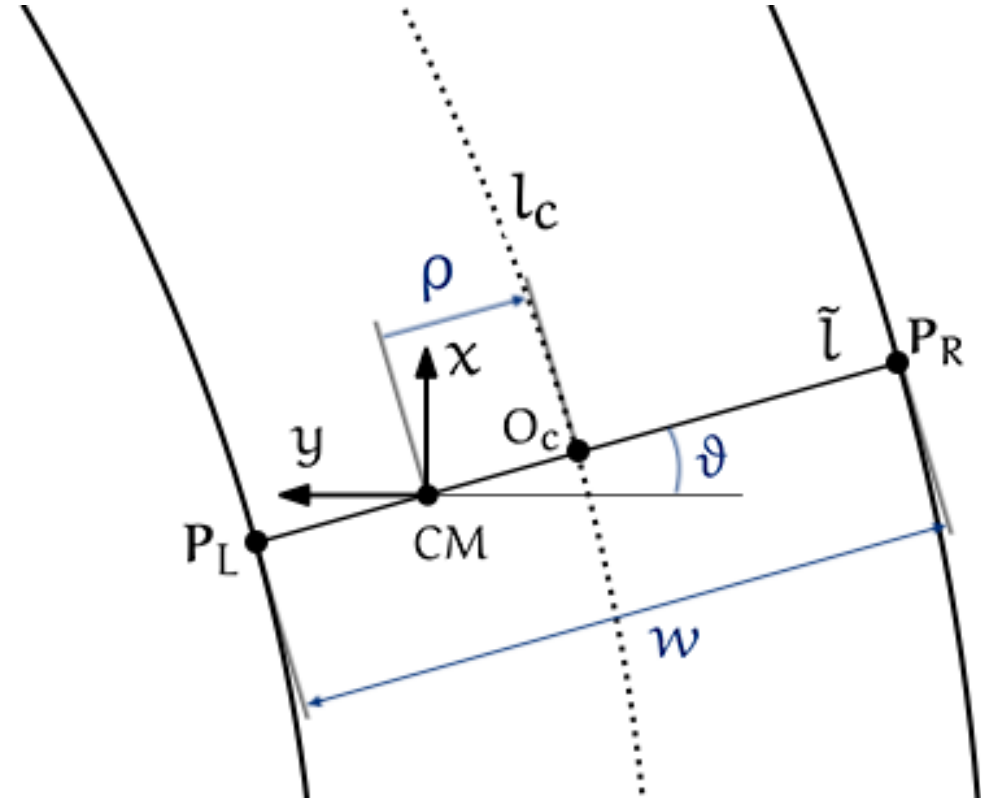


Ego vehicle relative position

Heading and lateral offset need to be computed exactly at CM

Requires the line $\tilde{l}$ such that:

- $CM \in \tilde{l}$
 - $\tilde{l} \perp l_c$
-
- From $\tilde{l}$ we can retrieve O_c and compute ρ
 - Temporal consistency guaranteed using an Extended Kalman Filter which considers (ϑ, ρ, w)



Test areas

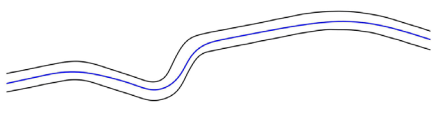
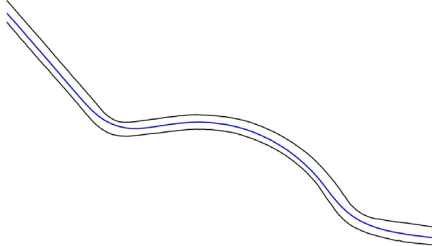
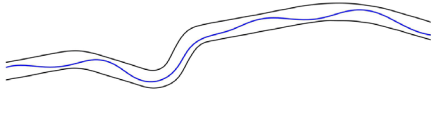
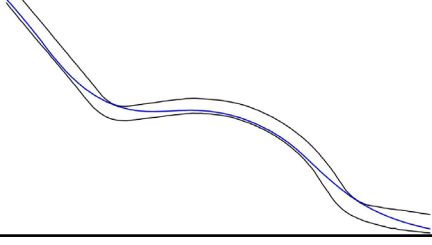
Monza ENI Circuit (5.4Km)



Arese test track (1.2Km)

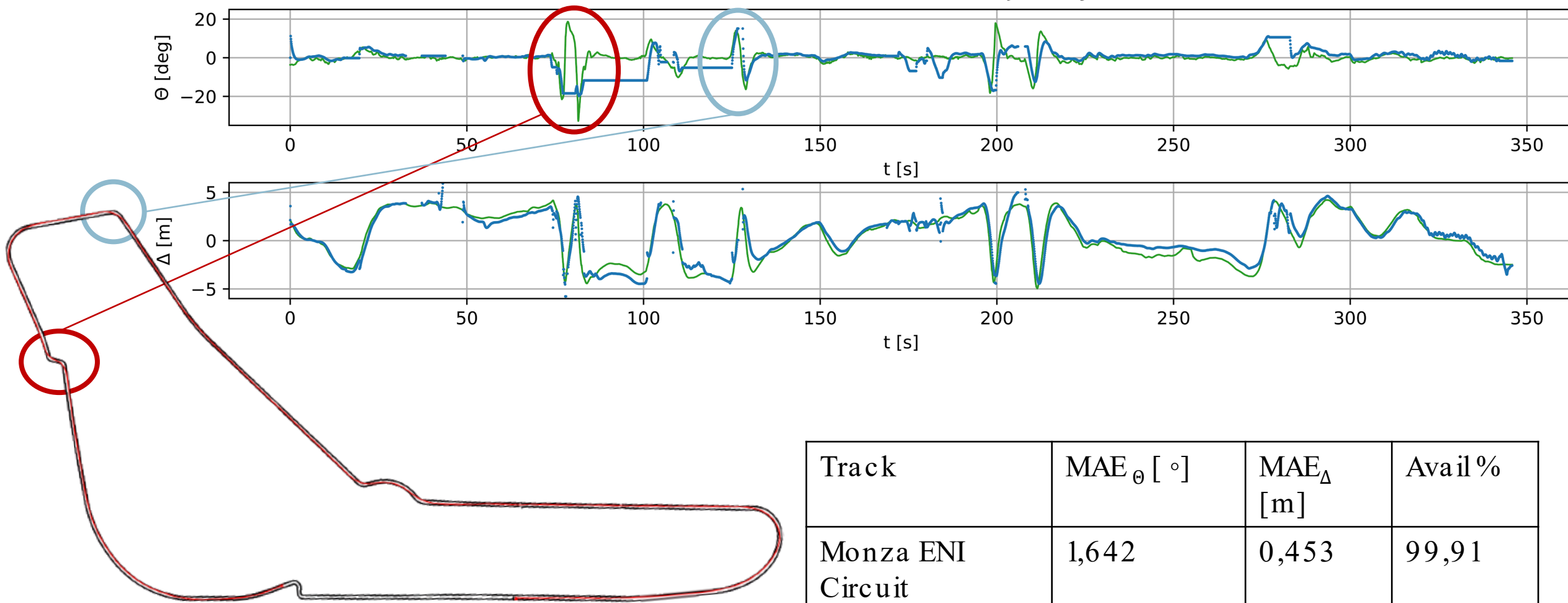


Ego vehicle relative position

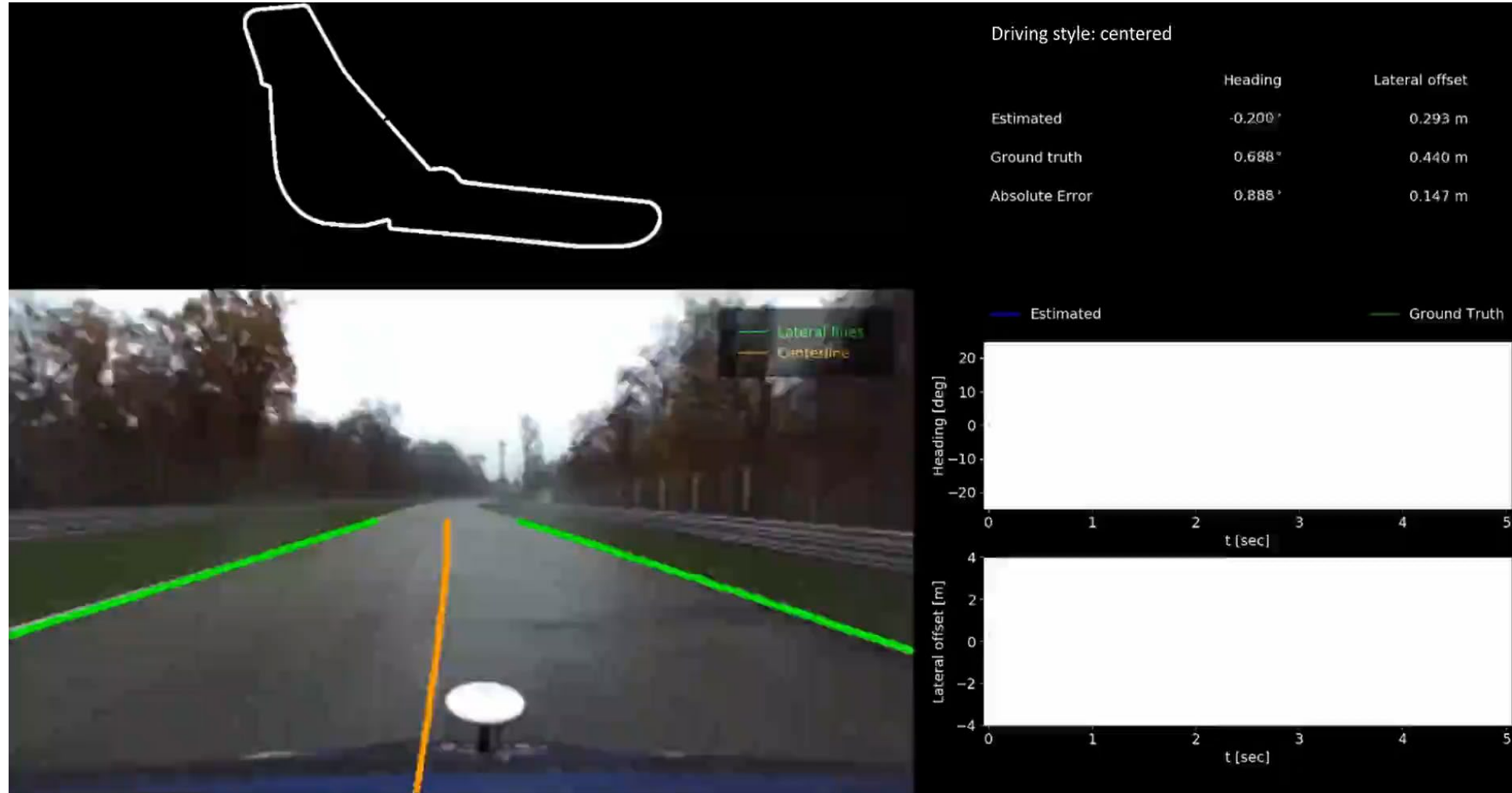
	Trajectory sample	$MAE_{\theta} [^{\circ}]$	$MAE_{\Delta} [m]$	Avail %
Centered track A (Arese Circuit)		2,892	0,820	99,71
Centered track B (Monza ENI Circuit)		1,642	0,453	99,91
Oscillating track A (Arese Circuit)		3,862	0,946	100
Racing track B (Monza ENI Circuit)		3,12	0,581	95,92

Ego vehicle relative position

Track B: Monza Eni Circuit - Trajectory: 2

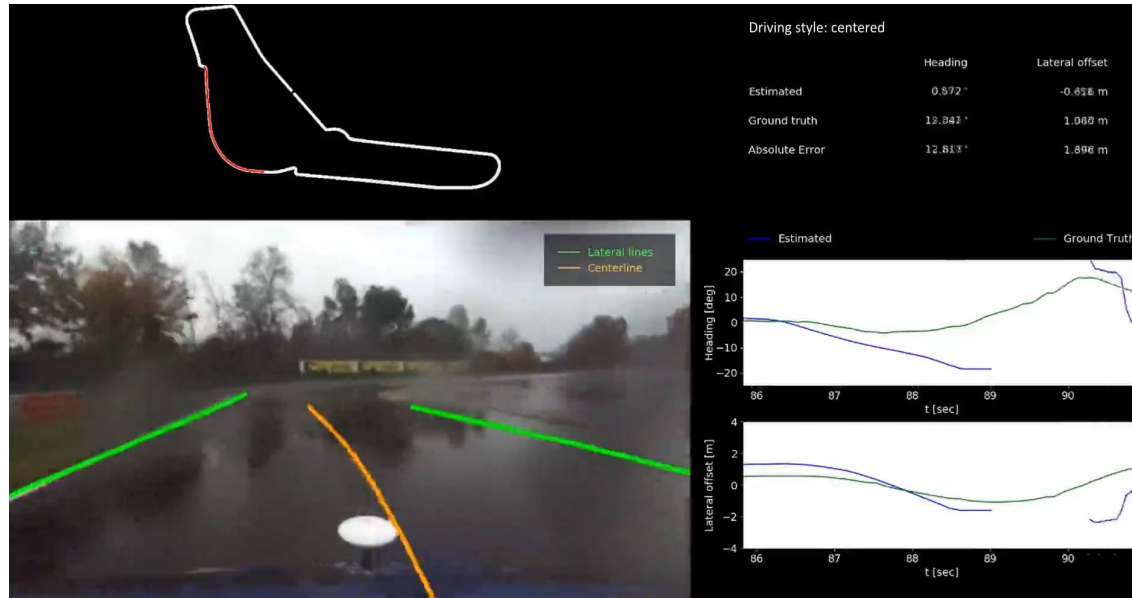


Vision -based ego vehicle state estimation

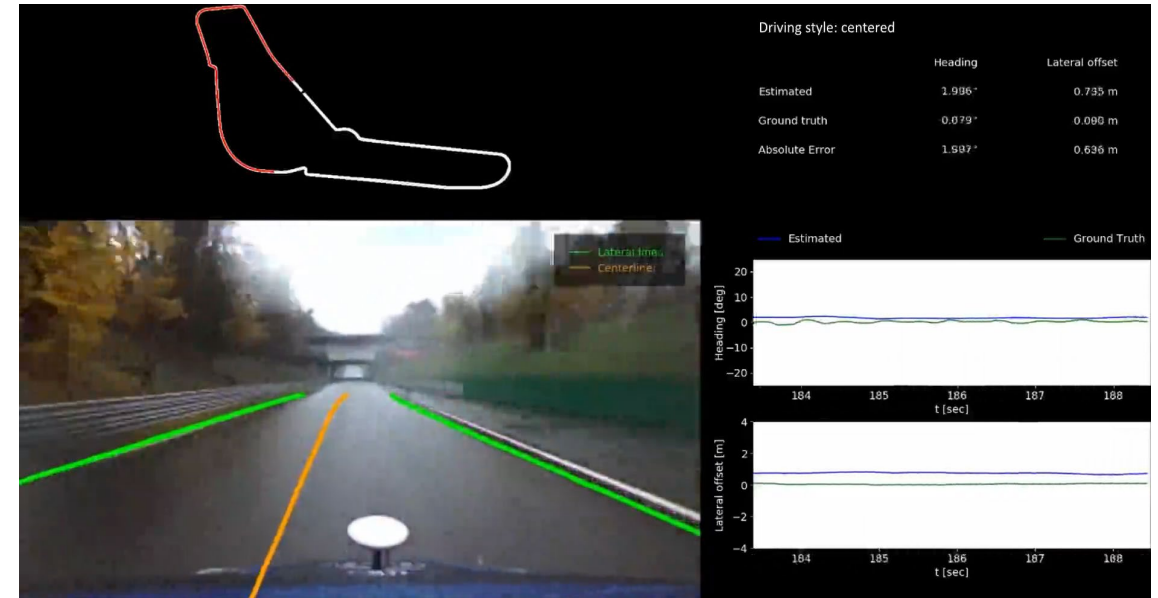


Challenging scenarios

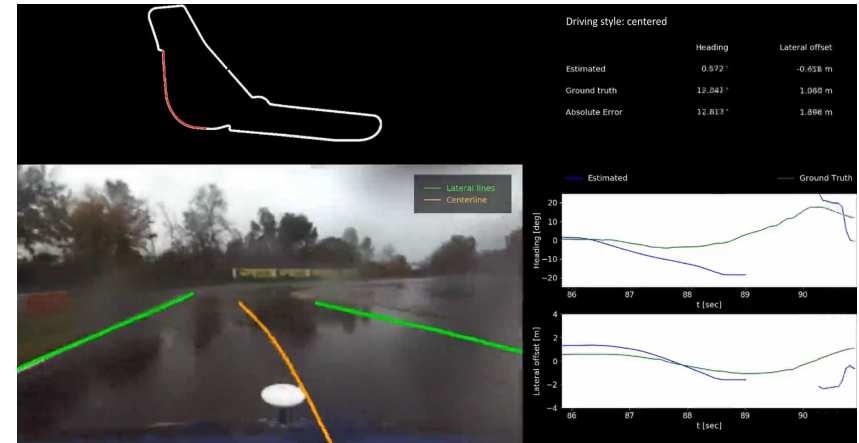
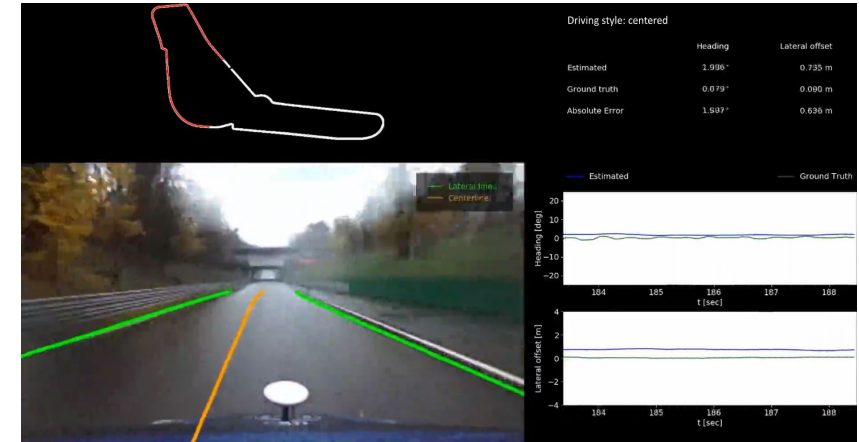
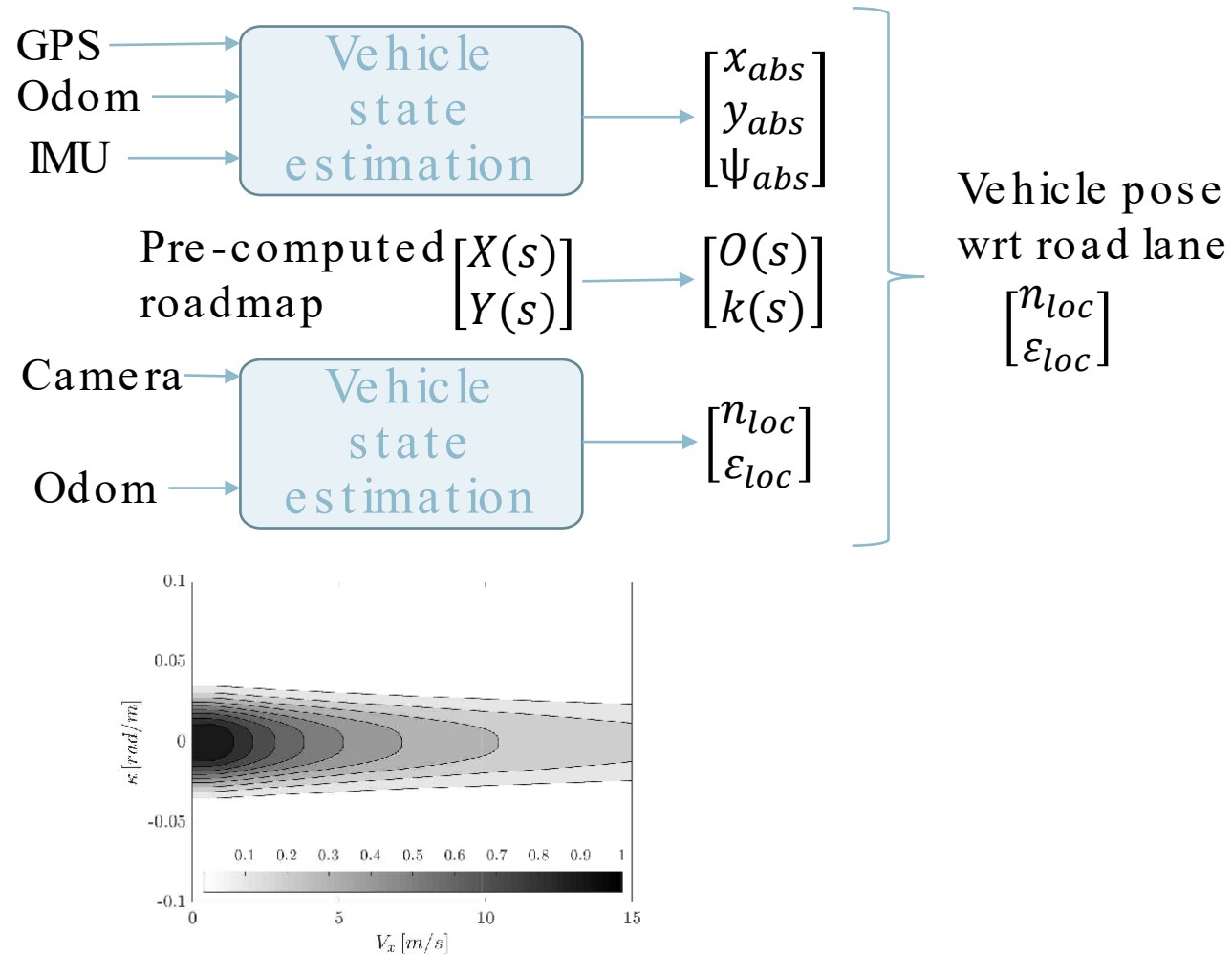
High curvature radius



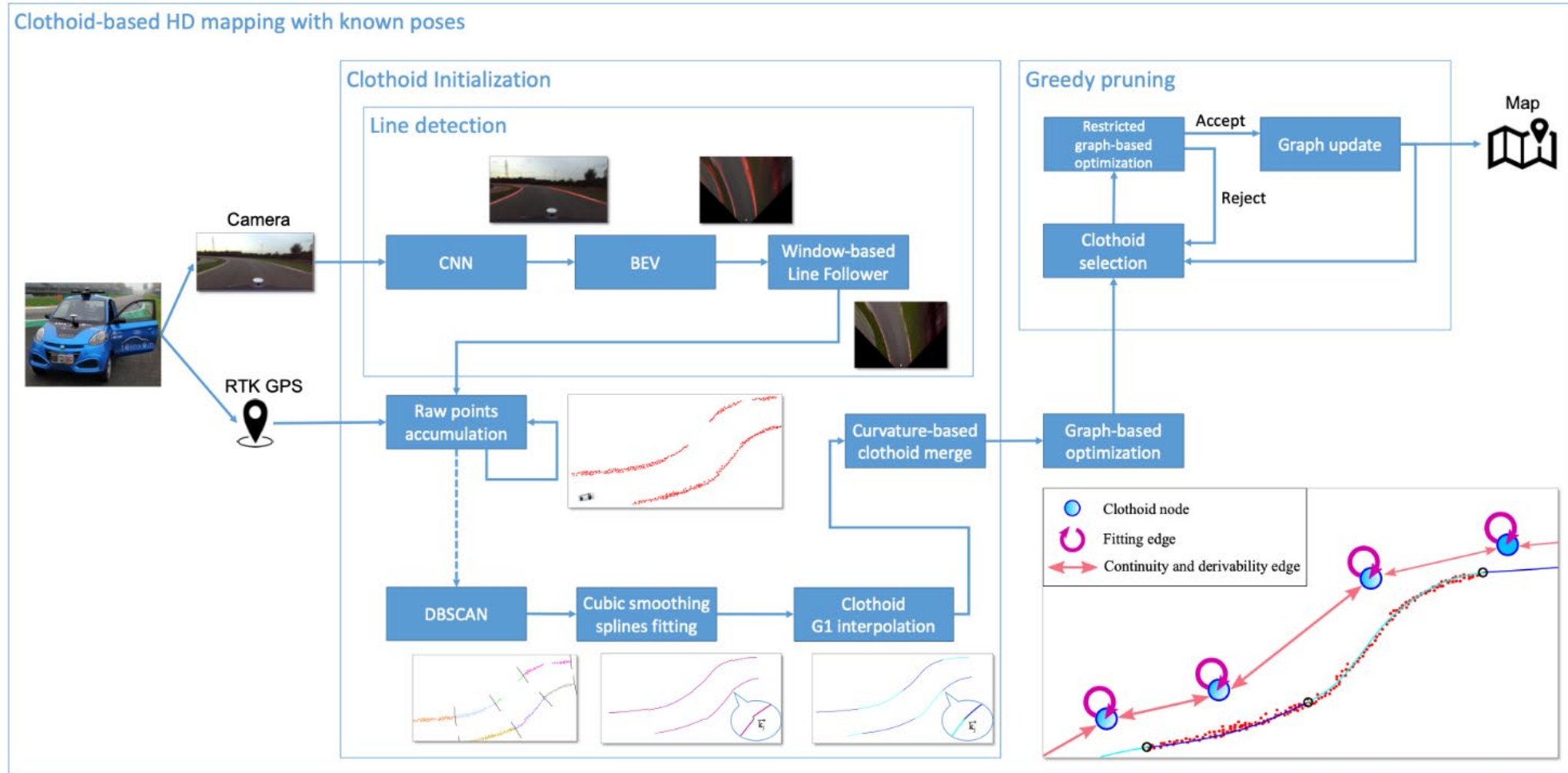
GNSS denied



Challenging scenarios



Offline road line detection - Mapping



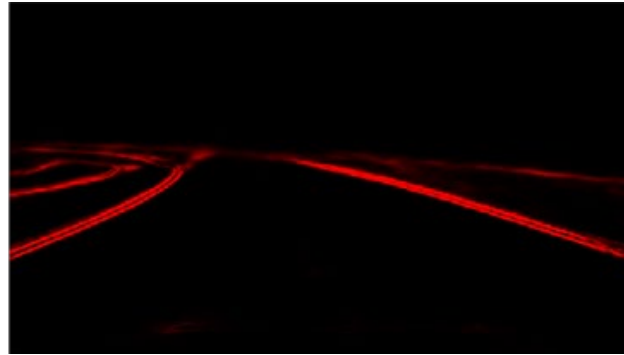
Mapping

Line Detection

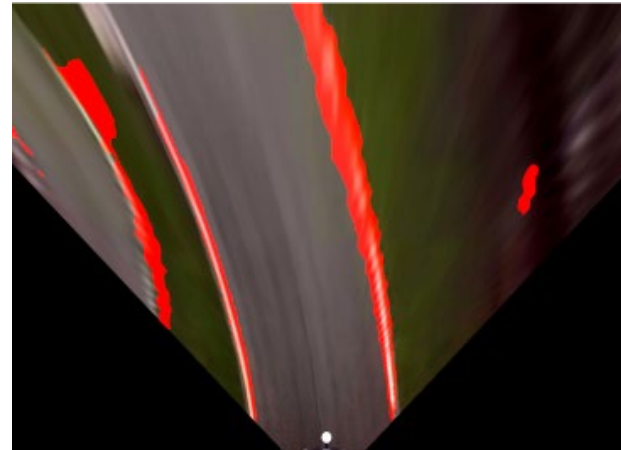
RGB



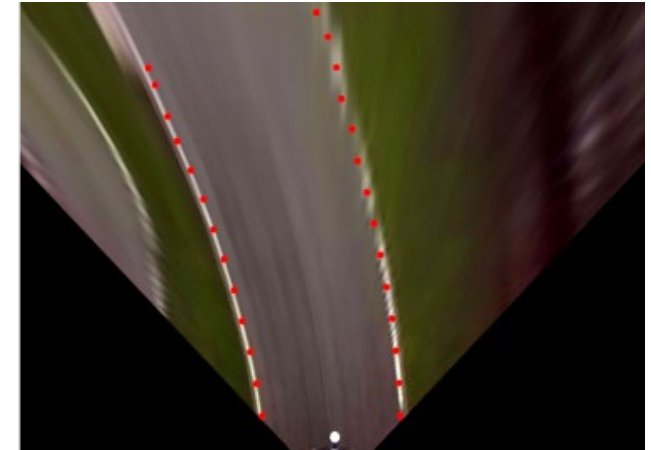
Segmentation



BEV



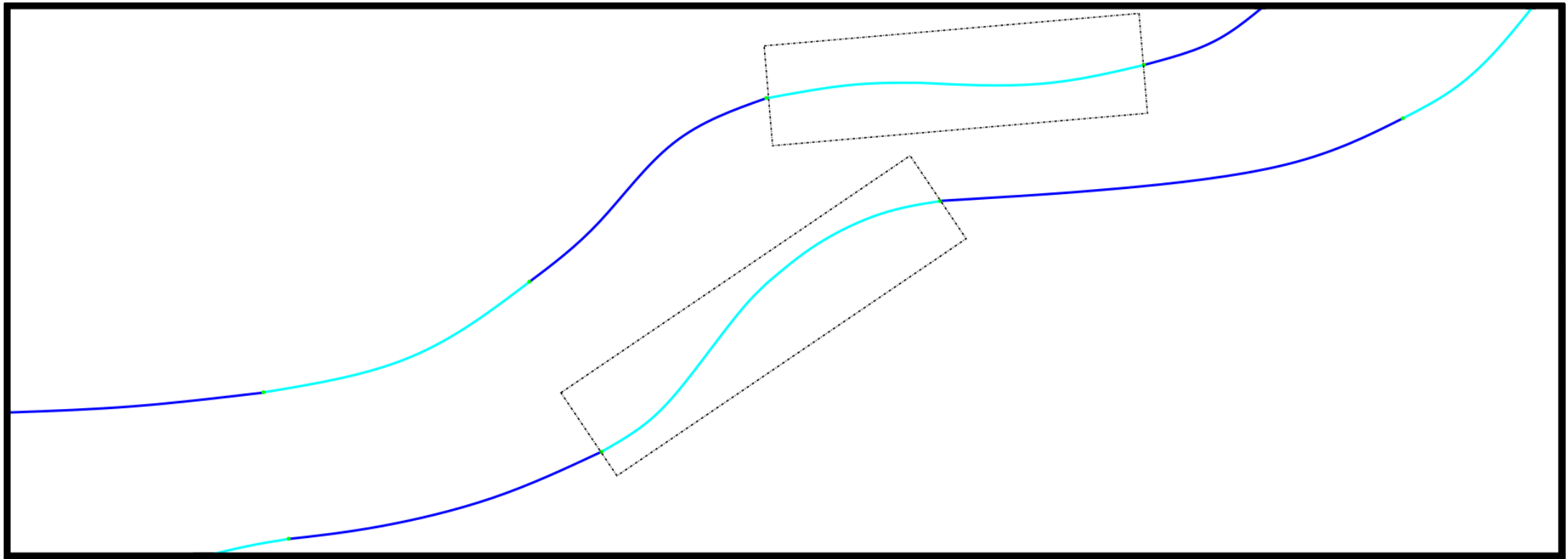
Points



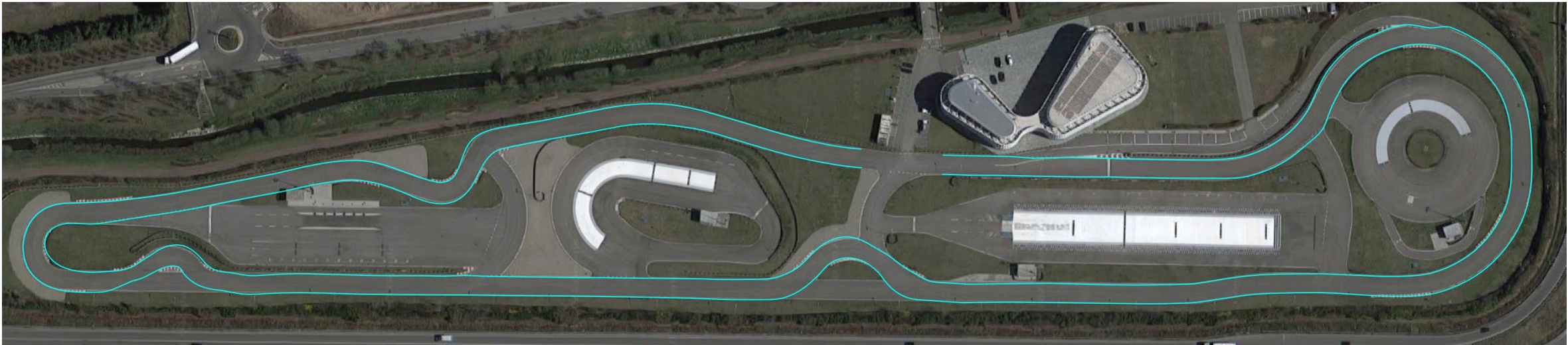
Mapping

Line Detection

Block division $\rightarrow$ DBSCAN $\rightarrow$ cubic smoothing splines $\rightarrow$ knots $\rightarrow$ clothoids $\rightarrow$ merging



Mapping – Arese



Clothoids quantity	MAE (m)	Continuity (m)		Derivability (°)	
		mean	max	mean	max
283 283	0.4103	0.0098	0.0113	0.1139	1.1931

Gallazzi, Barbara, Paolo Cudrano, Matteo Frosi, Simone Mentasti, and Matteo Matteucci. "Clothoidal mapping of road line markings for autonomous driving high-definition maps." In *2022 IEEE Intelligent Vehicles Symposium (IV)*, pp. 1631-1638. IEEE, 2022.
Cudrano, Paolo, Barbara Gallazzi, Matteo Frosi, Simone Mentasti, and Matteo Matteucci. "Clothoid-based lane-level high-definition maps: Unifying sensing and control models." *IEEE Vehicular Technology Magazine* 17, no. 4 (2022): 47-56.

Mapping – Monza



Clothoids quantity	MAE (m)	Continuity (m)		Derivability (°)	
		mean	max	mean	max
920	0.2772	0.0088	0.0103	0.0139	0.1064

Mapping – Monza

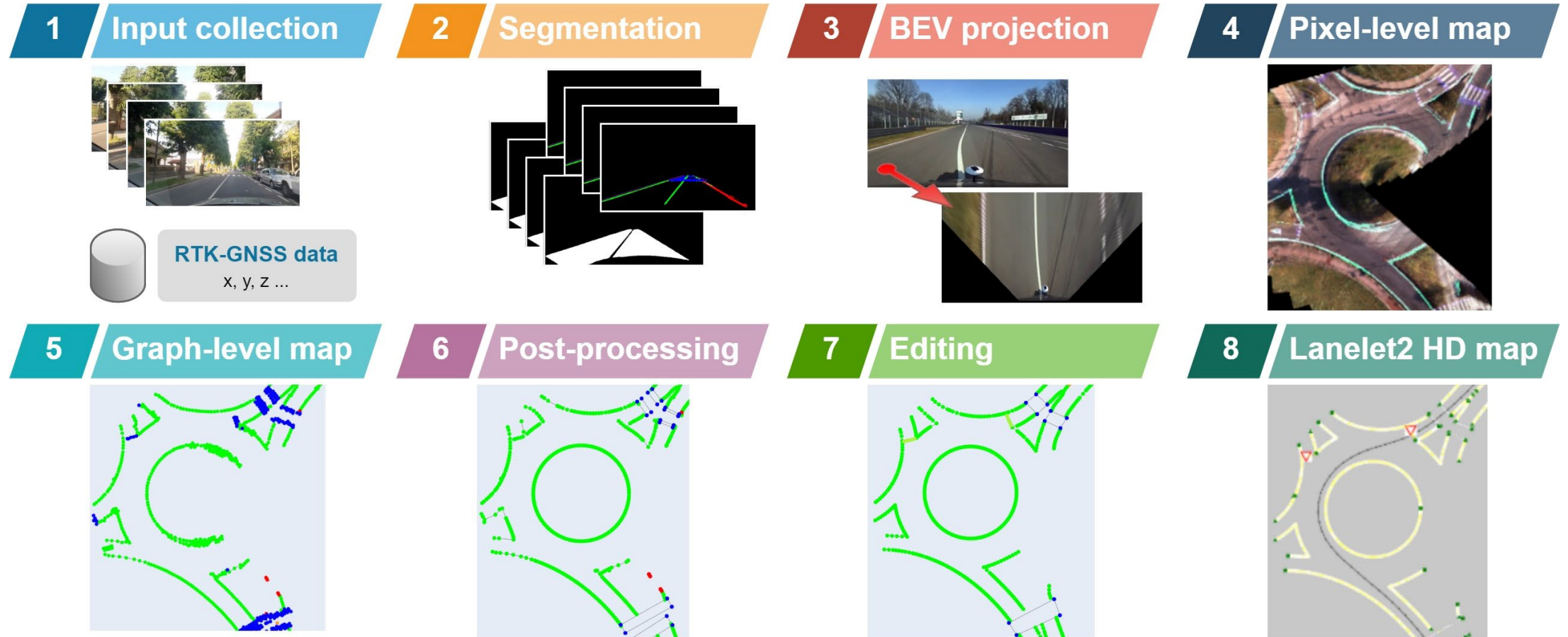


Gallazzi, Barbara, Paolo Cudrano, Matteo Frosi, Simone Mentasti, and Matteo Matteucci. "Clothoidal mapping of road line markings for autonomous driving high-definition maps." In *2022 IEEE Intelligent Vehicles Symposium (IV)*, pp. 1631-1638. IEEE, 2022.

Cudrano, Paolo, Barbara Gallazzi, Matteo Frosi, Simone Mentasti, and Matteo Matteucci. "Clothoid-based lane-level high-definition maps: Unifying sensing and control models." *IEEE Vehicular Technology Magazine* 17, no. 4 (2022): 47-56.

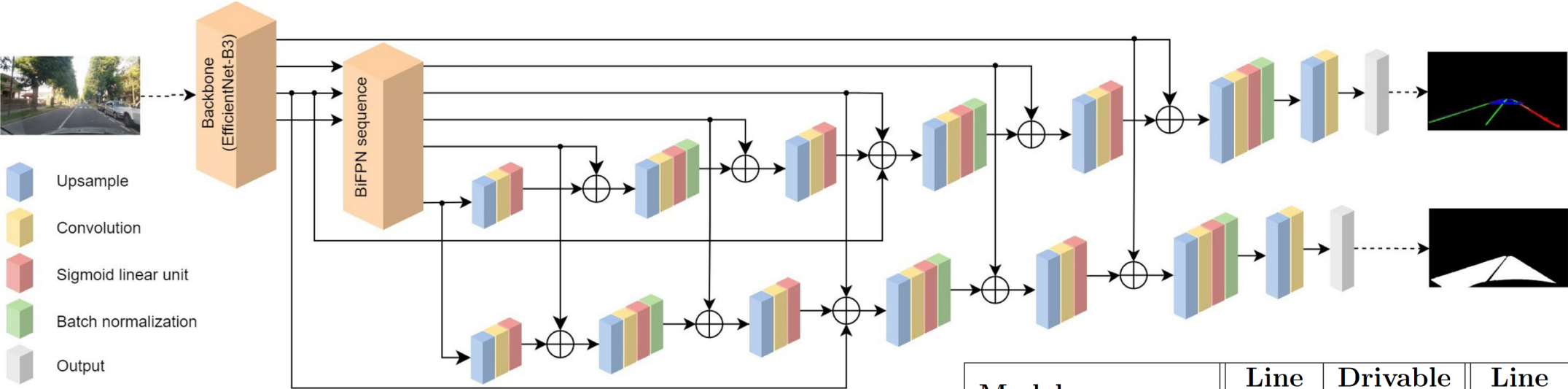
Semantic Mapping

On the road



Semantic Mapping

On the road



Model	Line IoU	Drivable IoU	Line Acc	Drivable Acc
HybridNets	53.82	86.59	59.81	86.87
Yolop	49.70	86.00	65.40	97.40
YolopV2	53.24	88.41	80.11	91.31
RoadStarNet-FT	53.45	87.66	66.50	87.89
RoadStarNet-F*	44.41	88.34	82.44	89.09

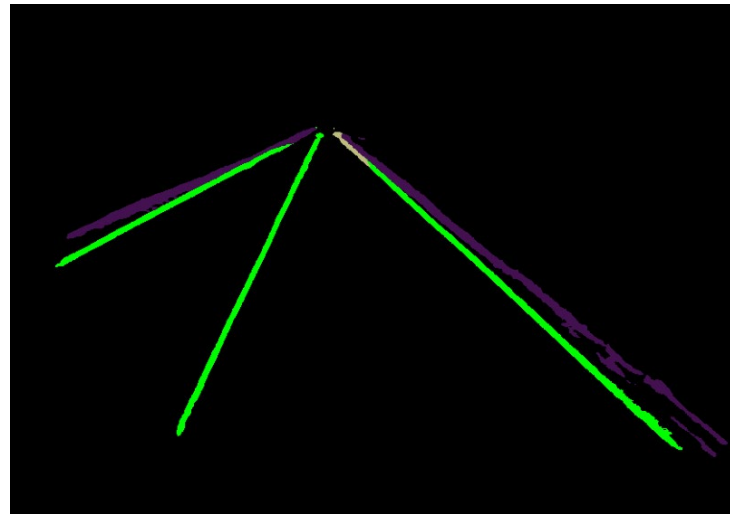
Semantic Mapping

On the road

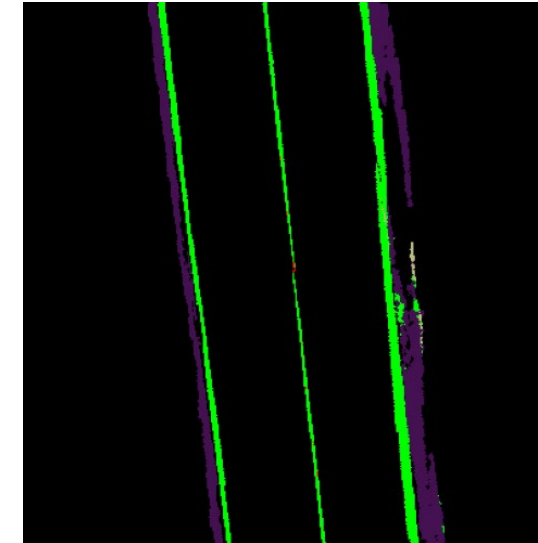
RGB image



Multiclass segmentation



BEV

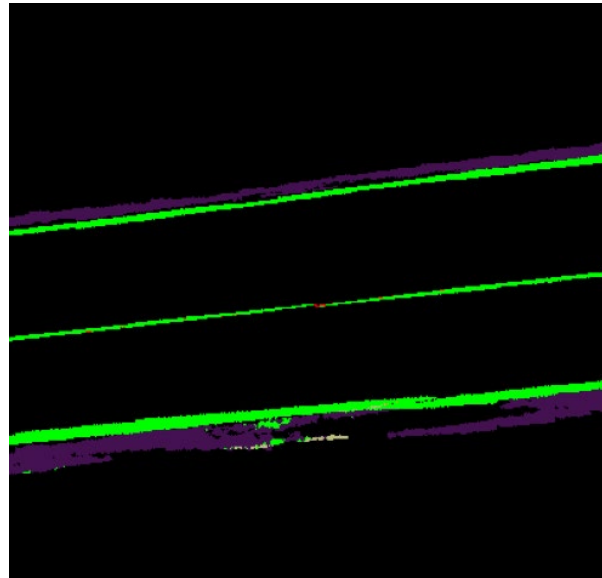


x: 362847.110, y: 5107972.381 z: 186.048
roll: 95.9881351, pitch: 89.7408898, yaw: 24.0095672

Semantic Mapping

On the road

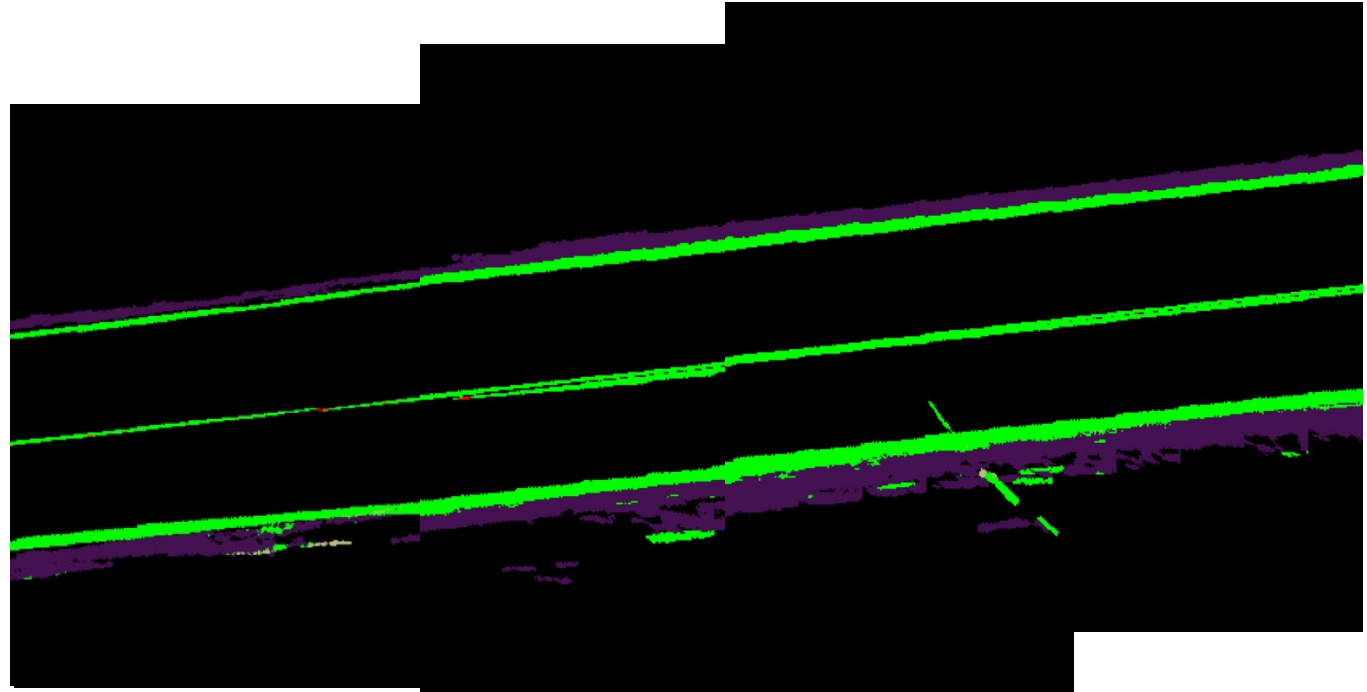
Aligned BEV using vehicle state



x: 362847.110, y: 5107972.381 z: 186.048

Semantic Mapping

On the road

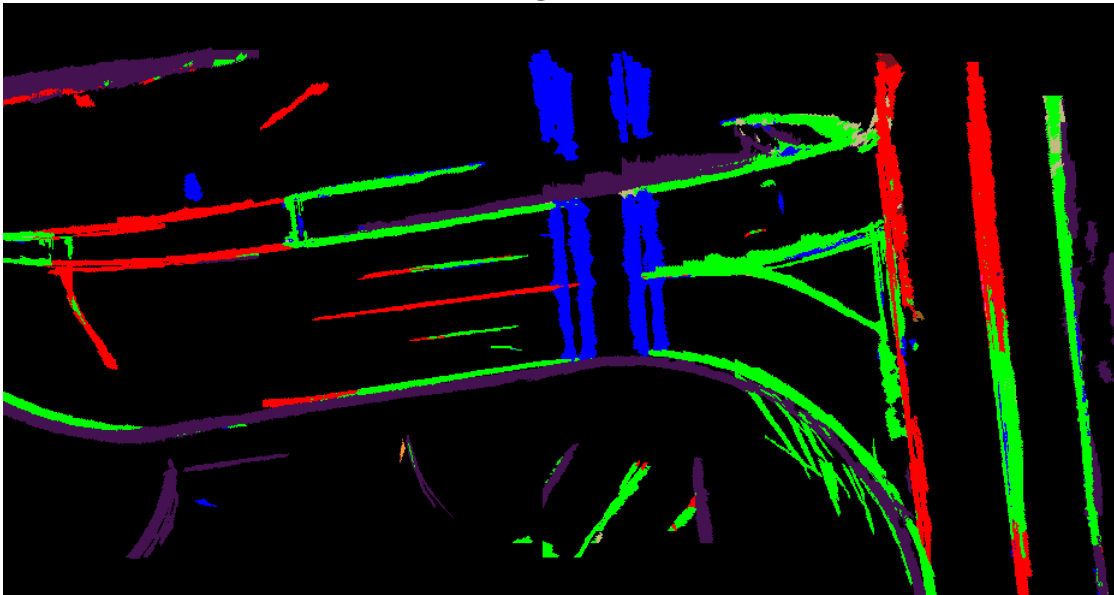


x: 362847.110, y: 5107972.381 z: 186.048

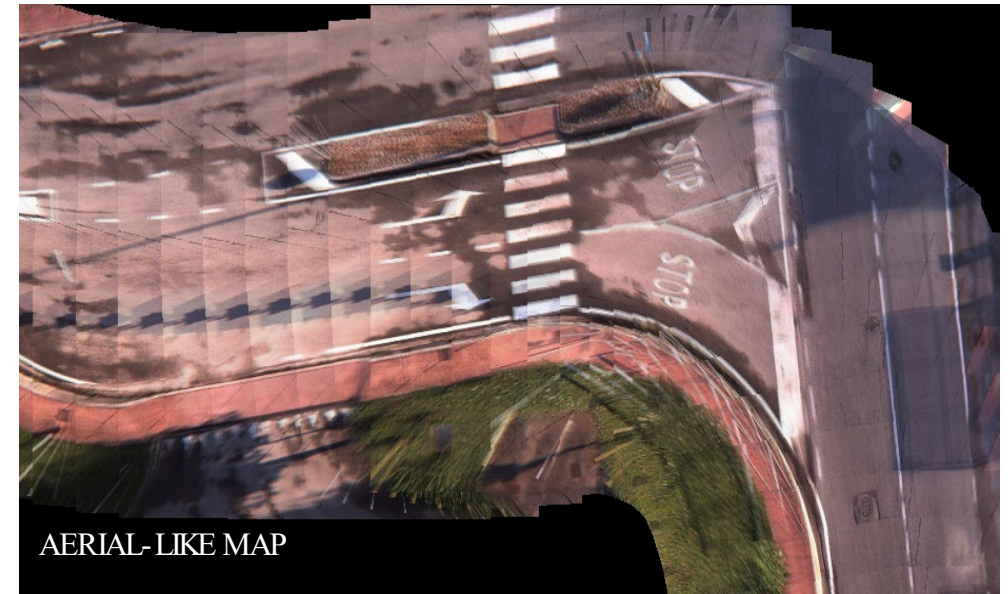
Semantic Mapping

On the road

Combined Segmentation BEV



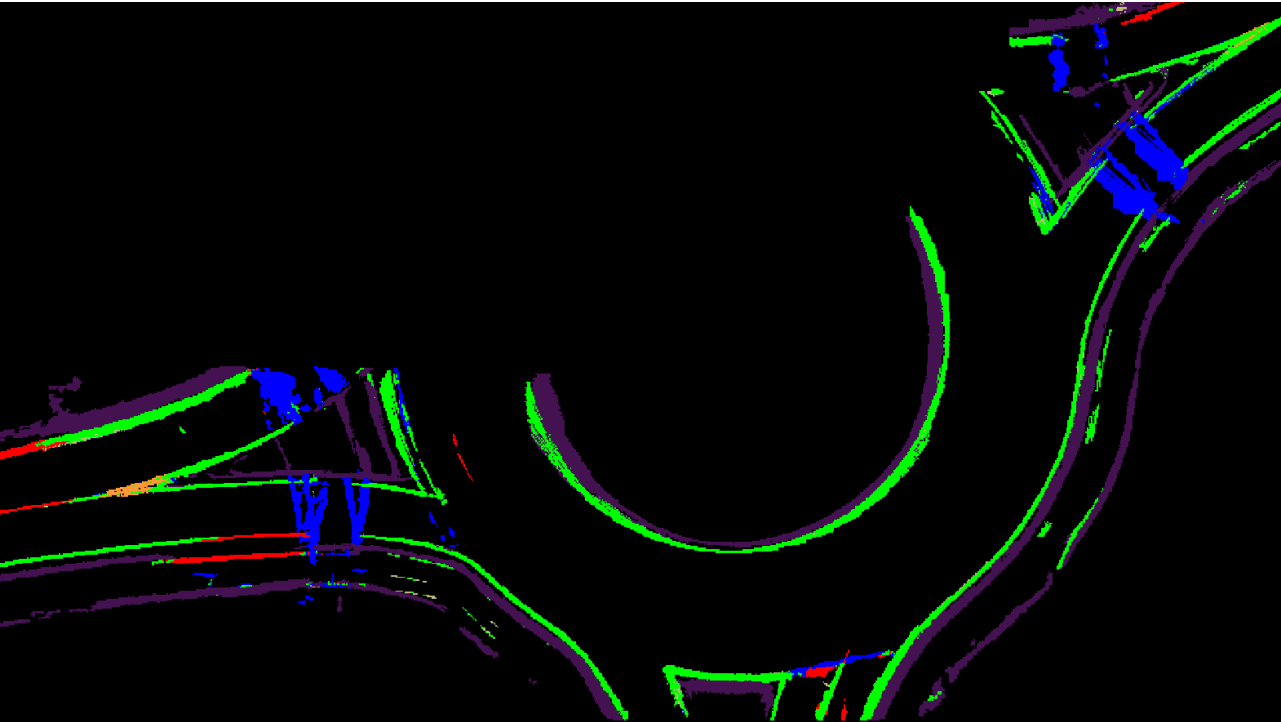
Combined RGB BEV



Semantic Mapping

On the road

Combined Segmentation BEV



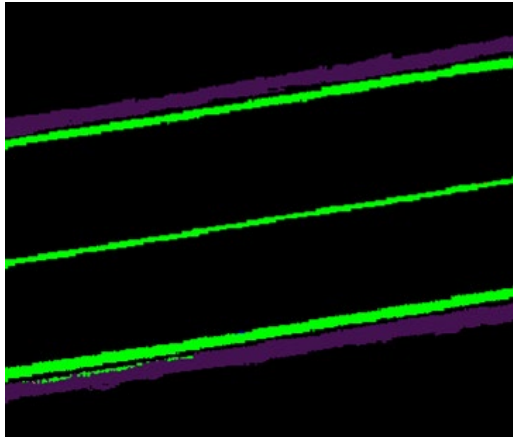
Monza Coverage

Model	Roggia	Lesmo	Ascari	Parabolica
HybridNets	0,3163m (85,49%)	0,3236m (61,69%)	0,3377m (64,76%)	0,3109m (64,70%)
Yolop	0,2933m (86,73%)	0,4222m (90,03%)	0,4162m (98,07%)	0,4074m (78,00%)
YolopV2	0,3723m (94,56%)	0,7340m (98,77%)	0,5955m (98,07%)	0,4545m (97,58%)
RoadStarNet-FT	0,2872m (83,64%)	0,4222m (83,69%)	0,3752m (90,20%)	0,3842m (80,30%)
RoadStarNet-F*	0,4047m (98,73%)	0,6277m (99,96%)	0,5457m (99,82%)	0,4750m (98,90%)

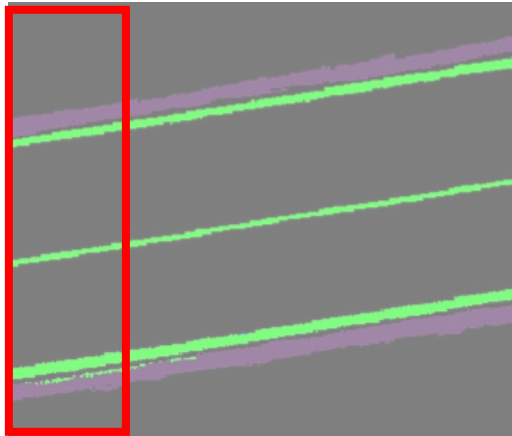
Semantic Mapping

Graph generation

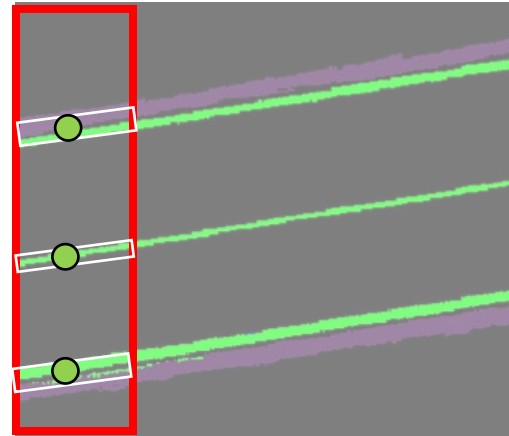
BEV



Fixed-size road area



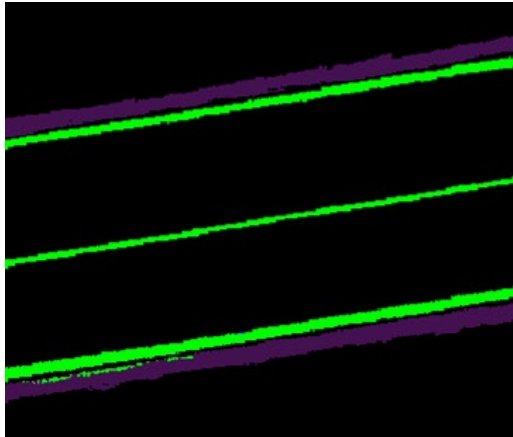
Line points



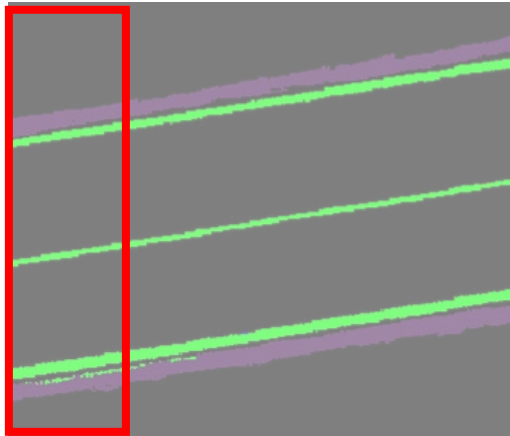
Semantic Mapping

Graph generation

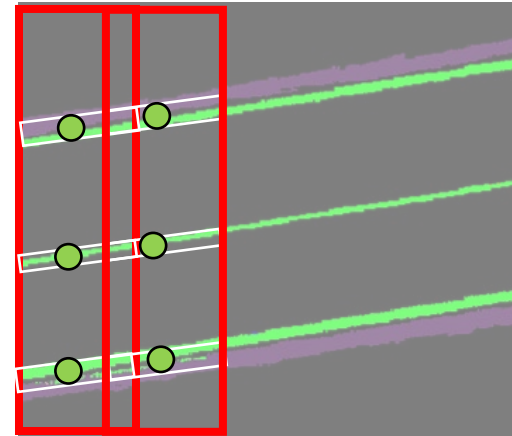
BEV



Fixed-size road area



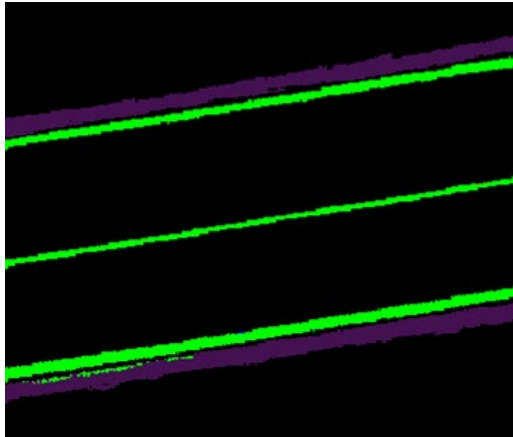
Sliding window



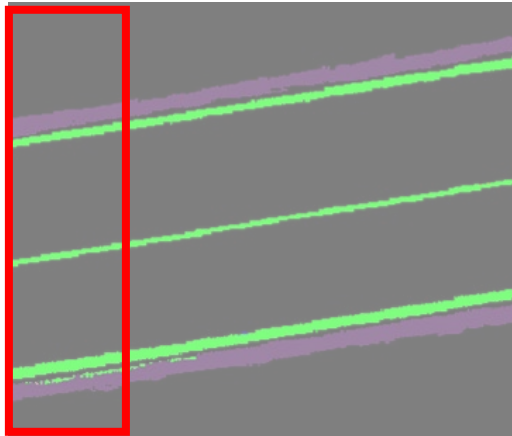
Semantic Mapping

Graph generation

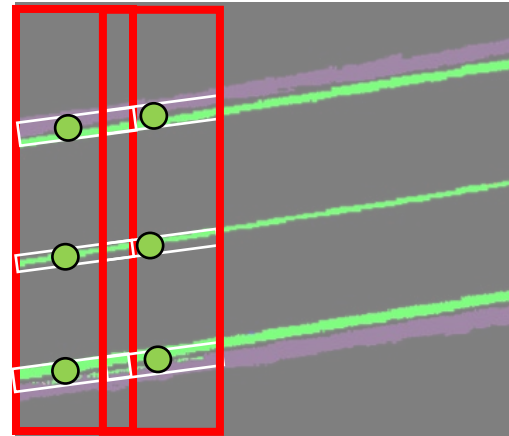
BEV



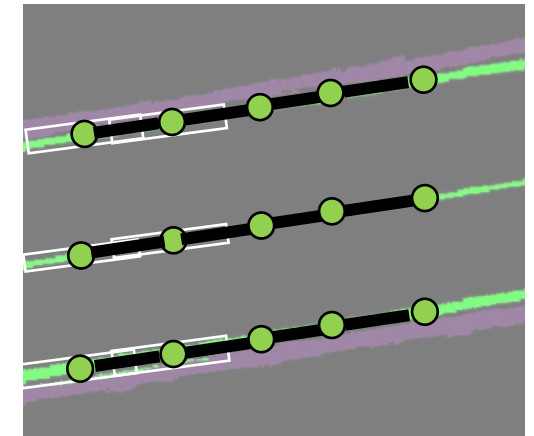
Fixed-size road area



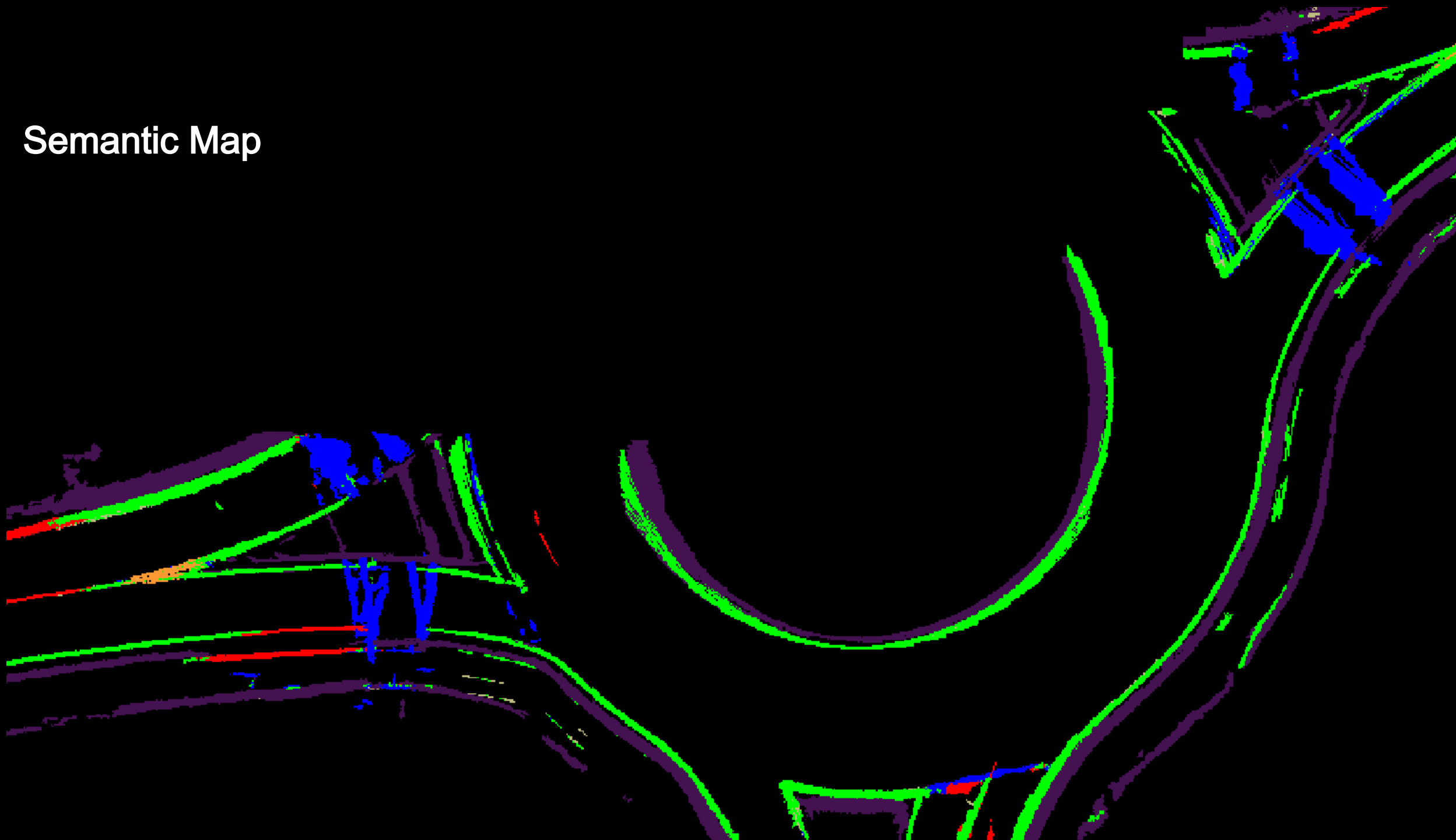
Sliding window



Line points



Semantic Map



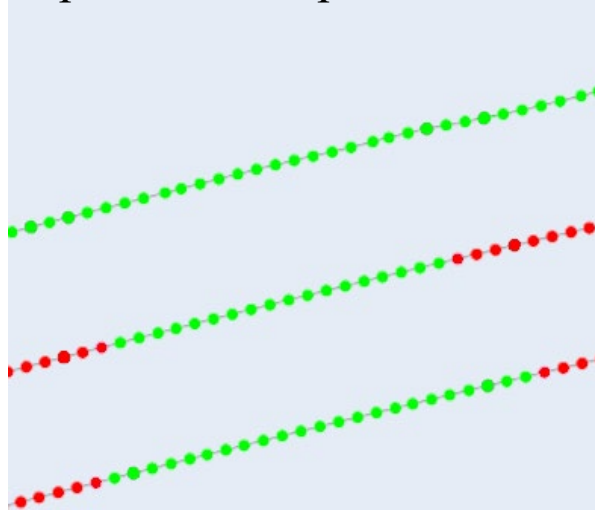
Graph-based map



Semantic Mapping

Graph reduction

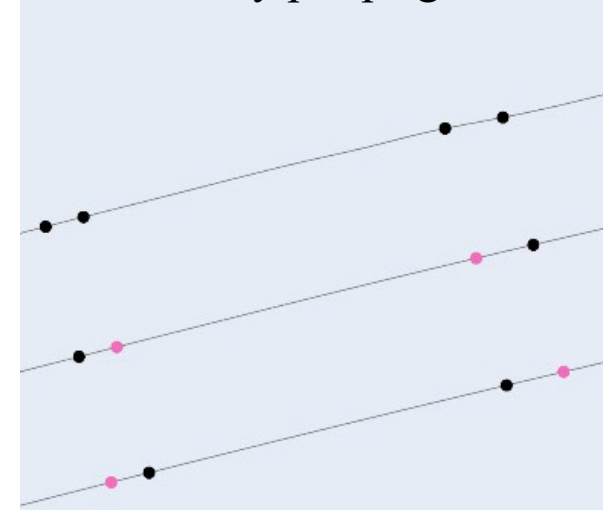
Graph based representation



Boundary detection



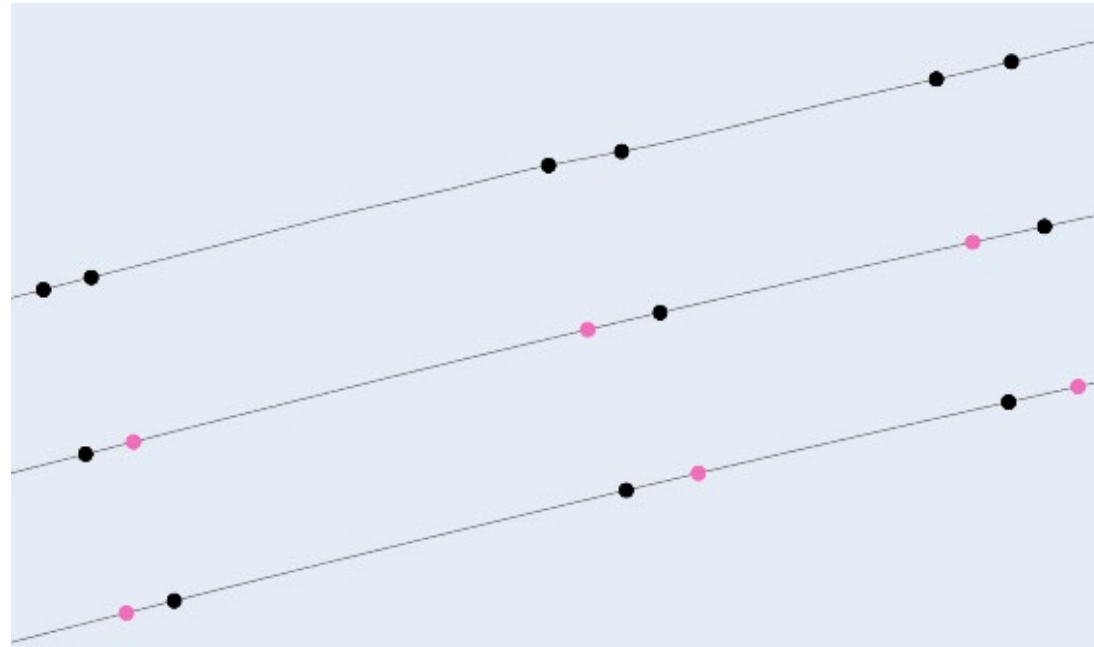
Boundary propagation



Semantic Mapping

Graph reduction

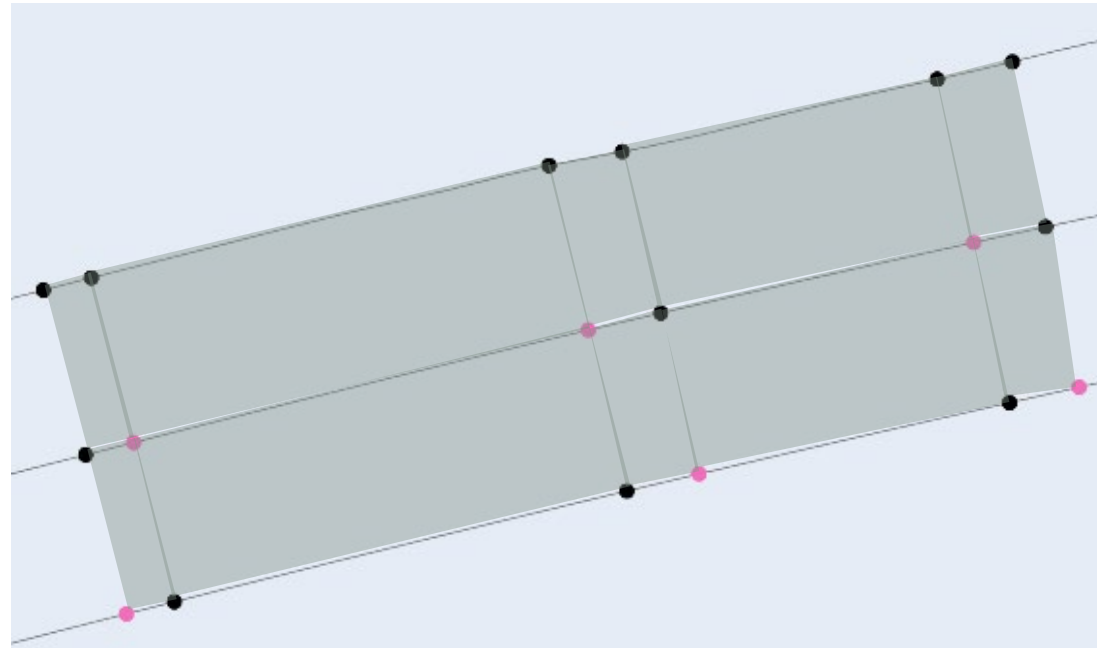
Minimal representation



Semantic Mapping

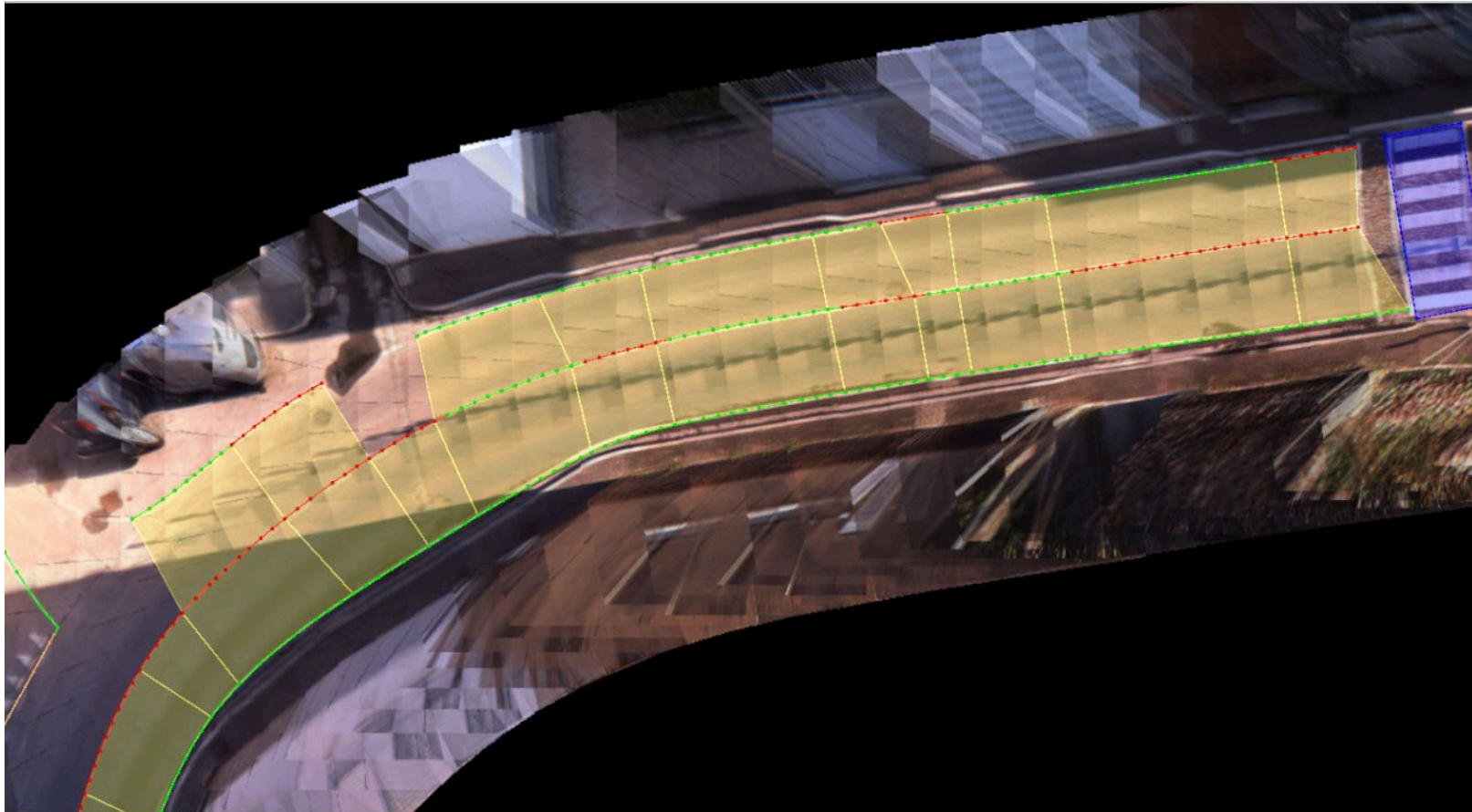
Lane Classification

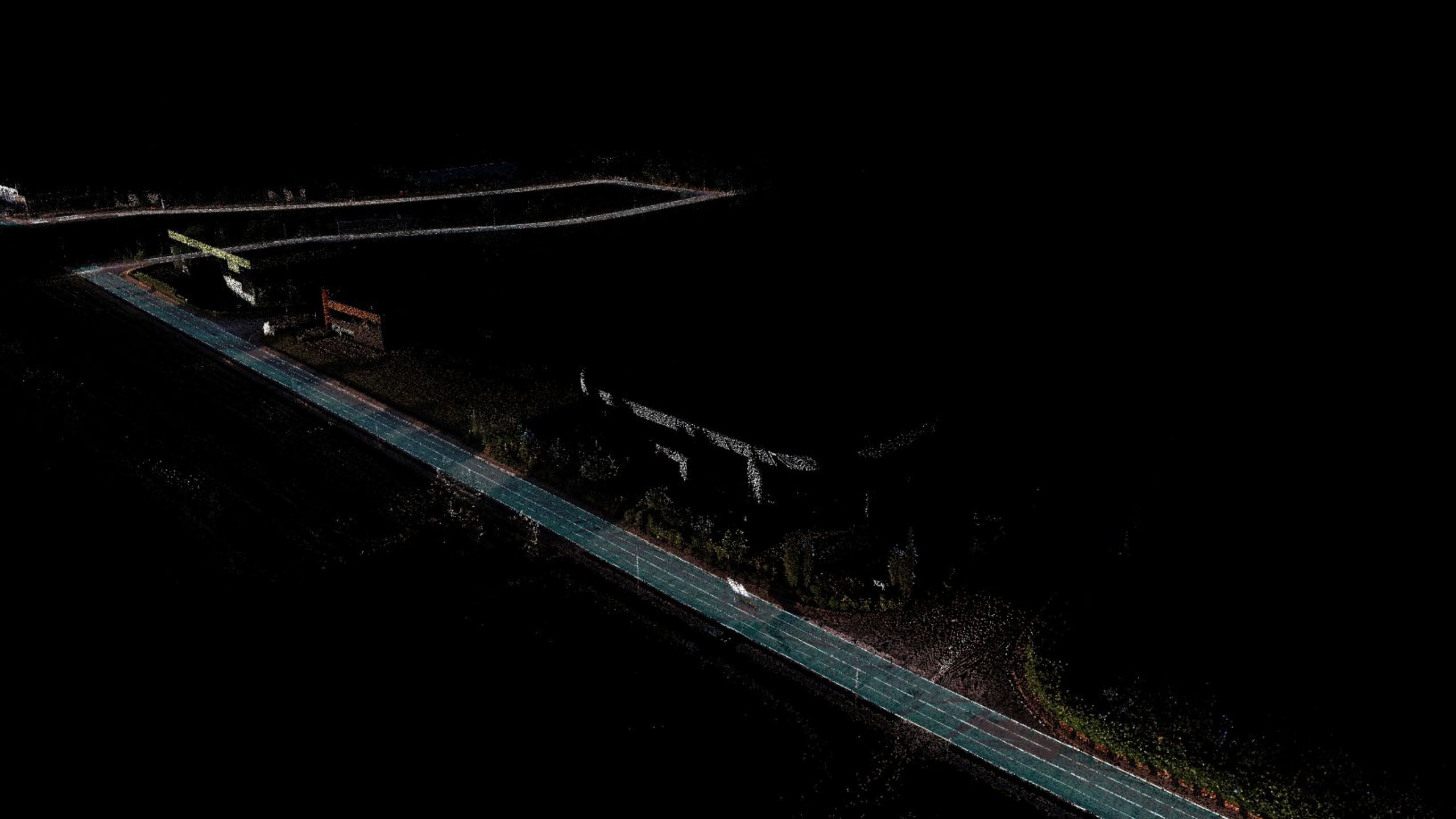
Lane boundaries definition



Semantic Mapping

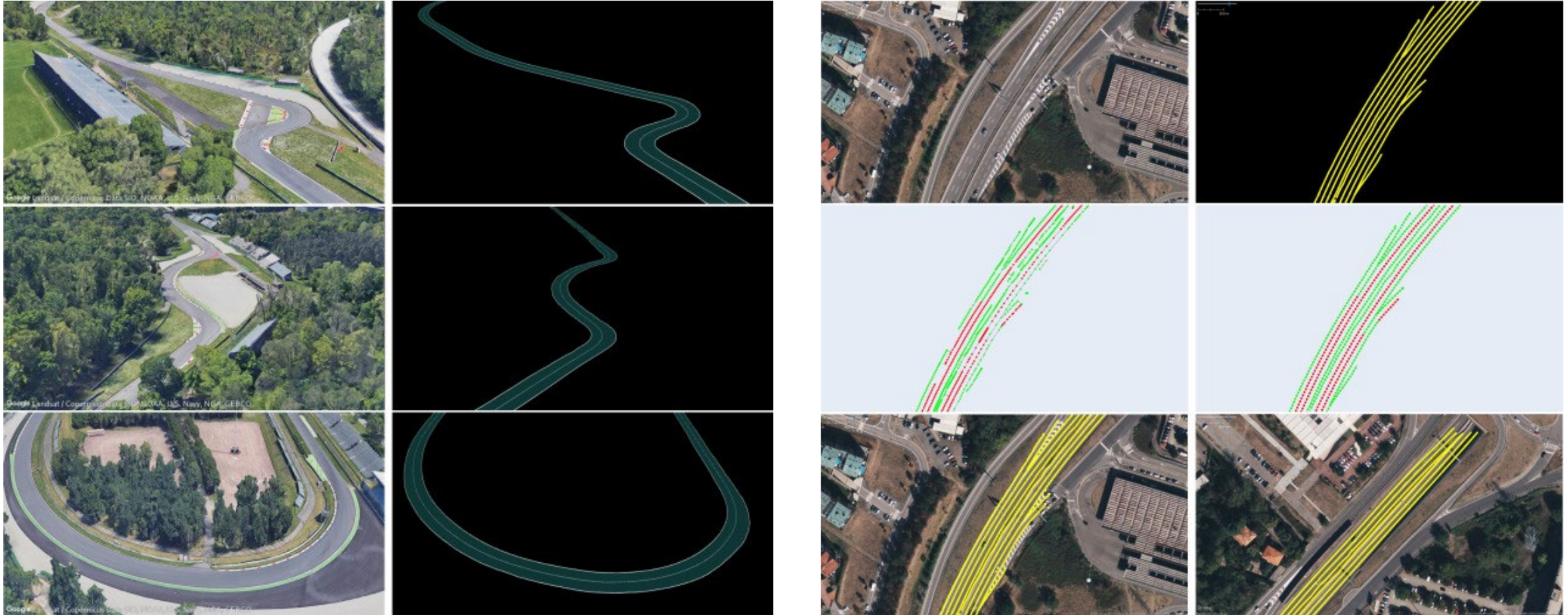
Lane Classification





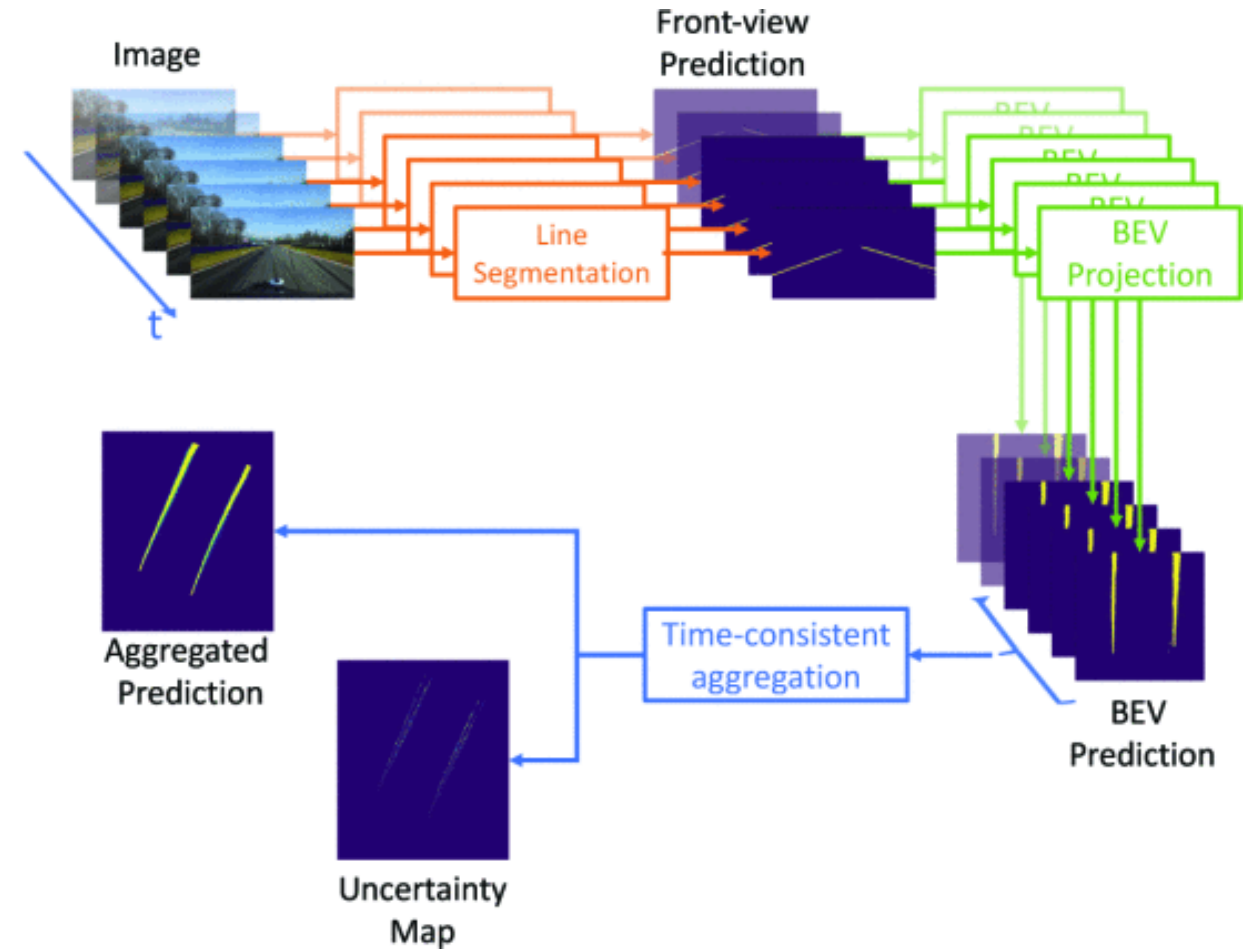
Semantic Mapping

Results



Temporal Consistency

And The Rise Of Transformers



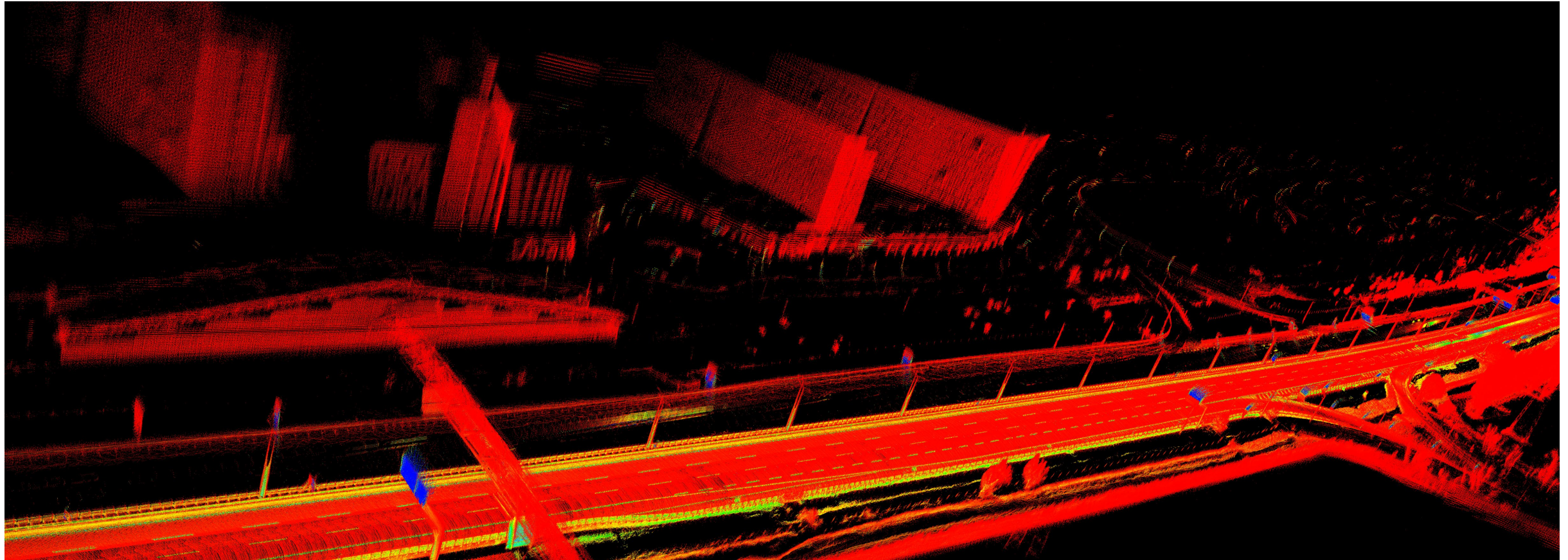
Why use cameras?

The power of high resolution LiDARs



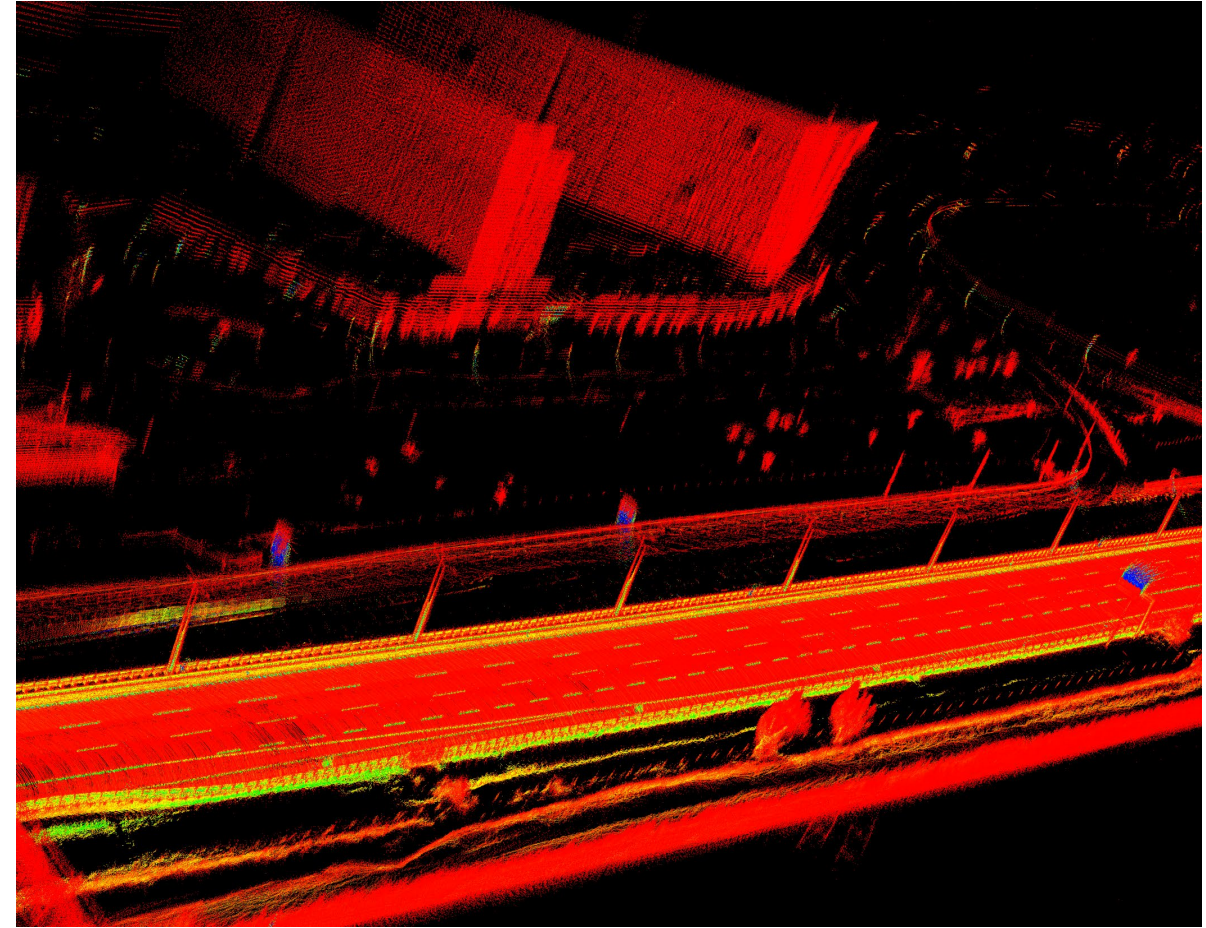
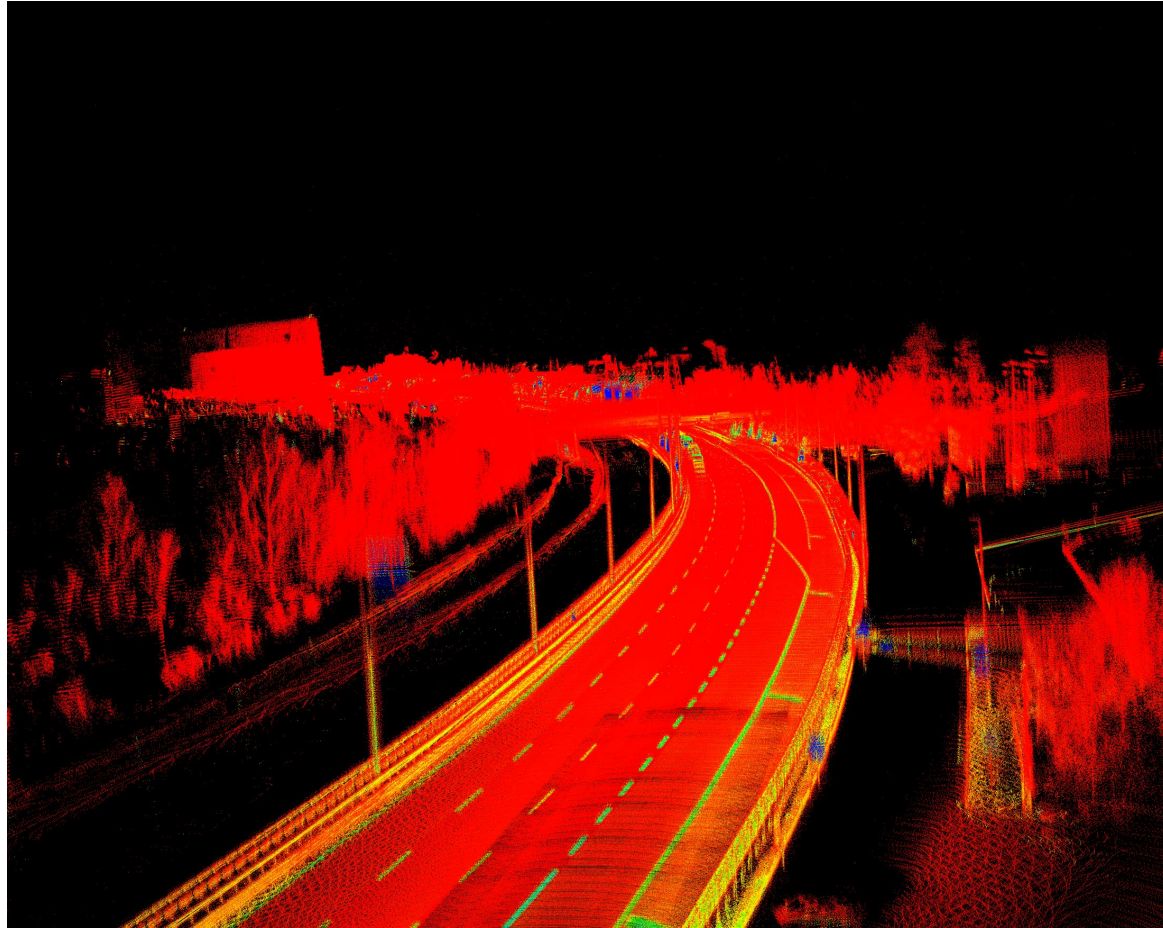
Why use cameras?

The power of high resolution LiDARs



Why use cameras?

The power of high resolution LiDARs



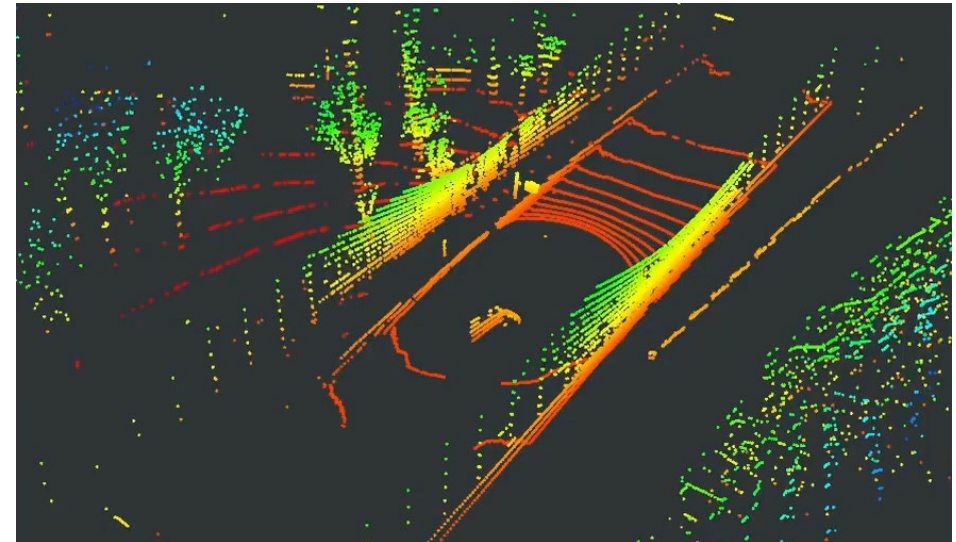
Seeing the others

Sensors with semantic

Cameras

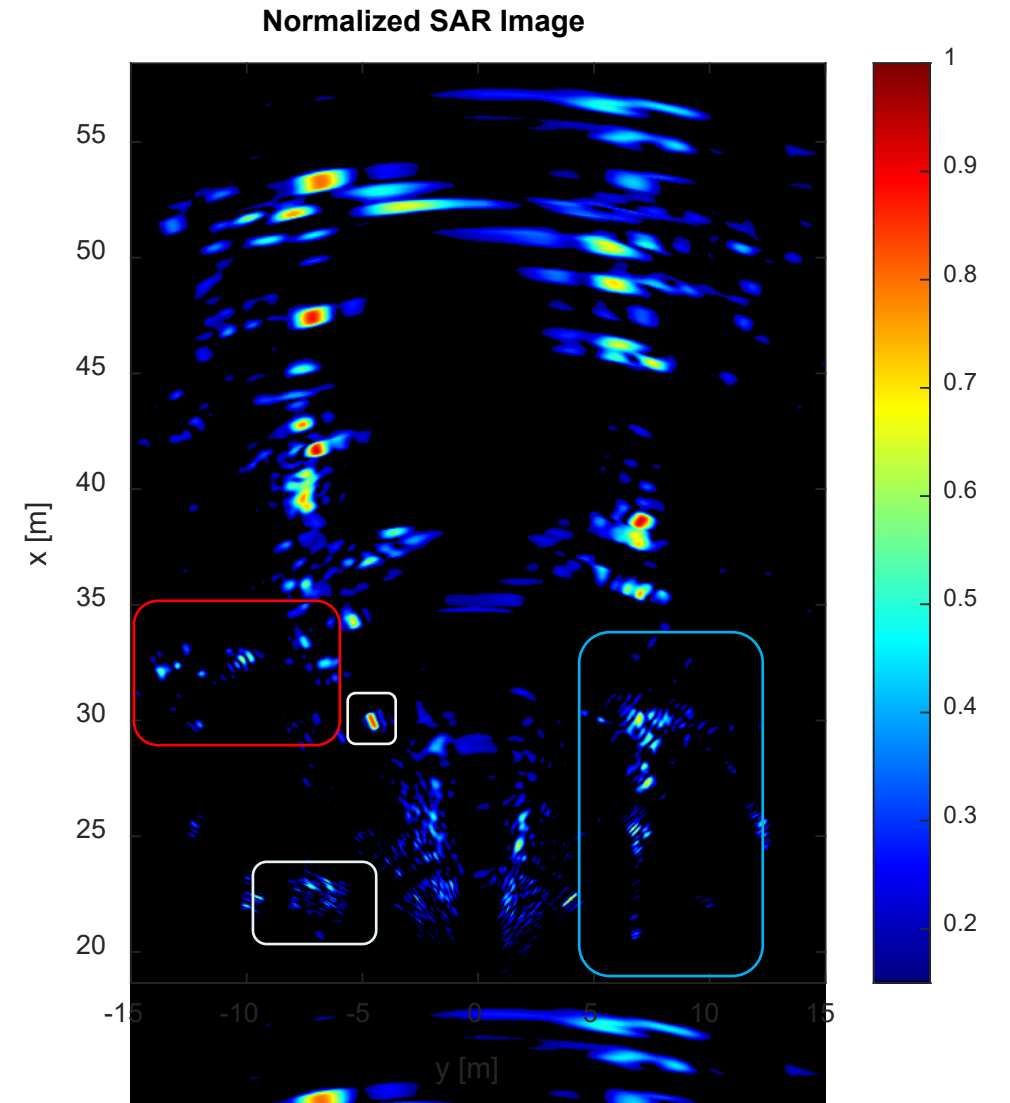
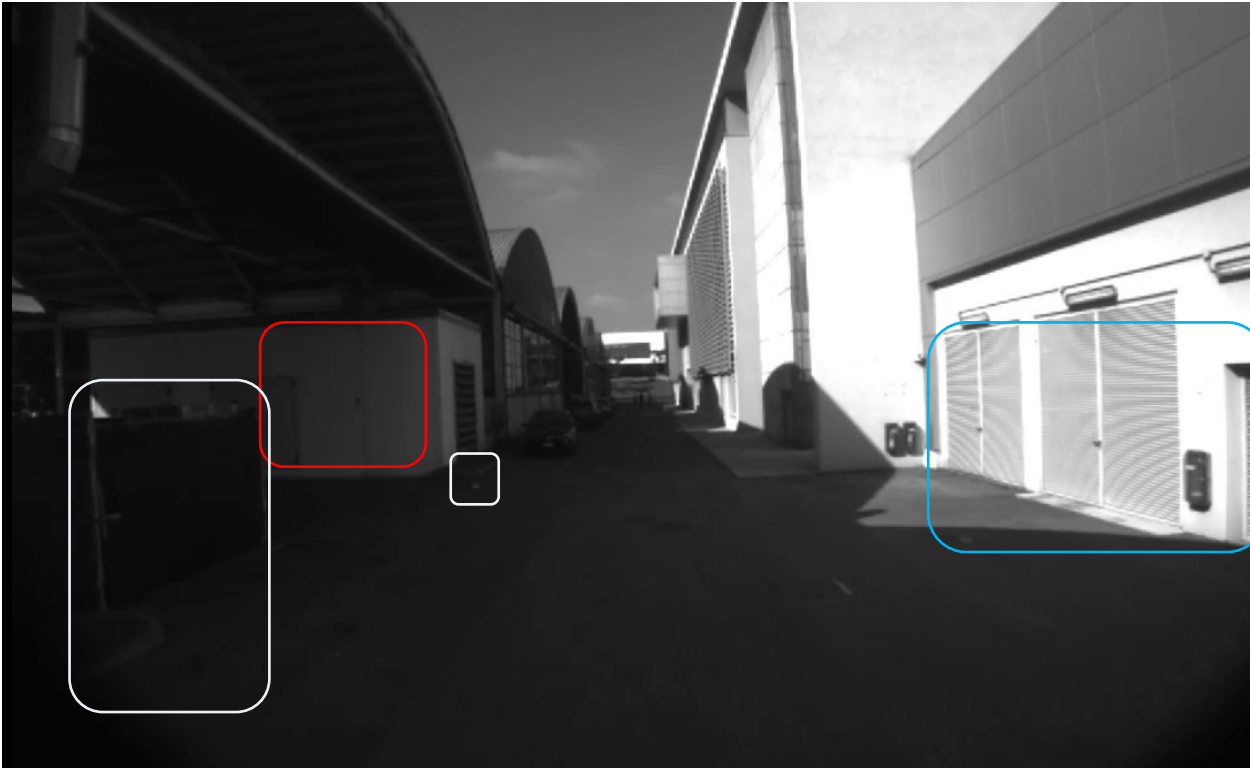


LiDARs



Seeing the others

Radars

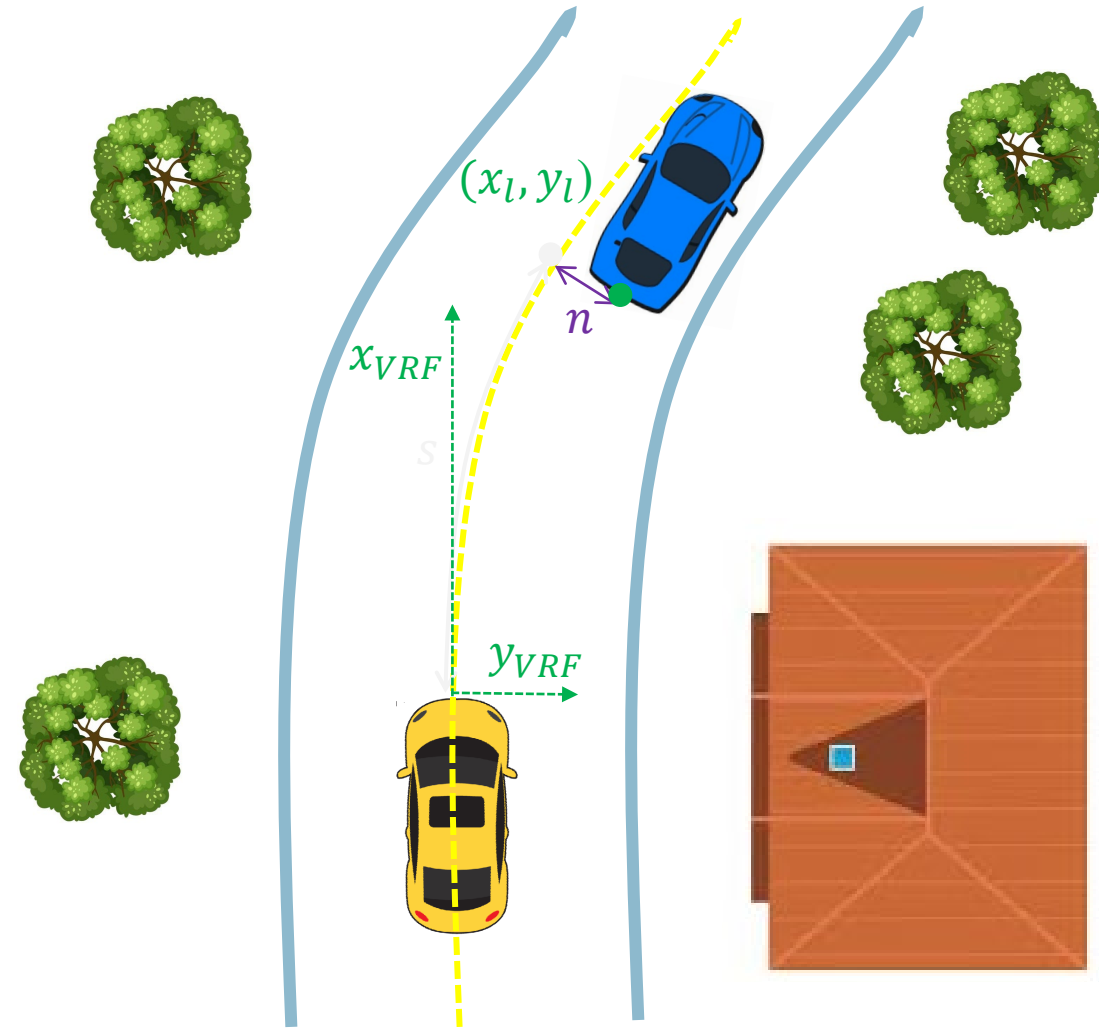


Seeing the others

Only the relevant ones

The object you are really interested in
are the one you are going to crash with

Represent object in trajectory space



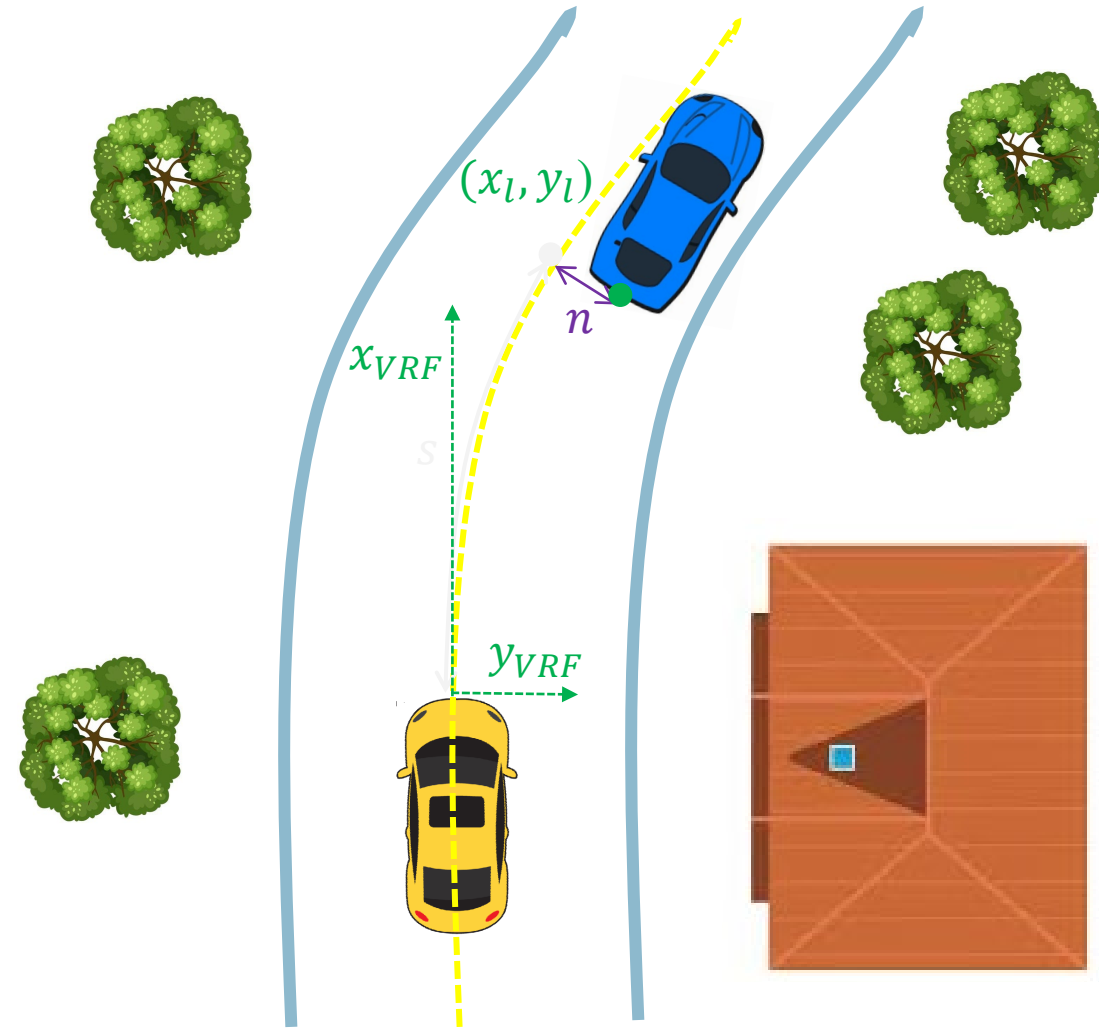
Seeing the others

In the correct reference frame

The object you are really interested in are the one you are going to crash with

Represent object in trajectory space

Sensors work in ego vehicle reference frame



Seeing the others

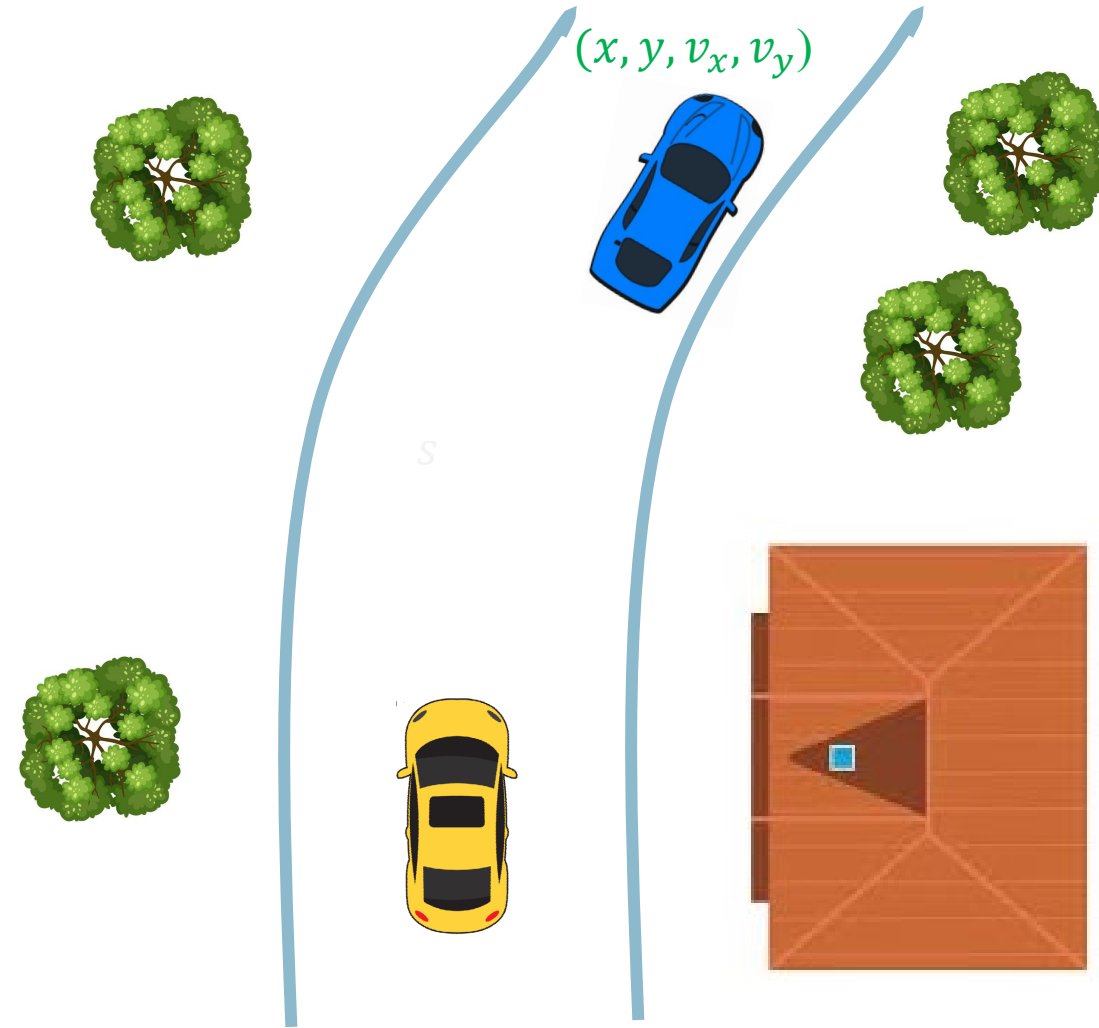
Radar

The object you are really interested in are the one you are going to crash with

Represent object in trajectory space

Sensors work in ego vehicle reference frame

RADAR (x, y, v_x, v_y)



Seeing the others

LiDAR

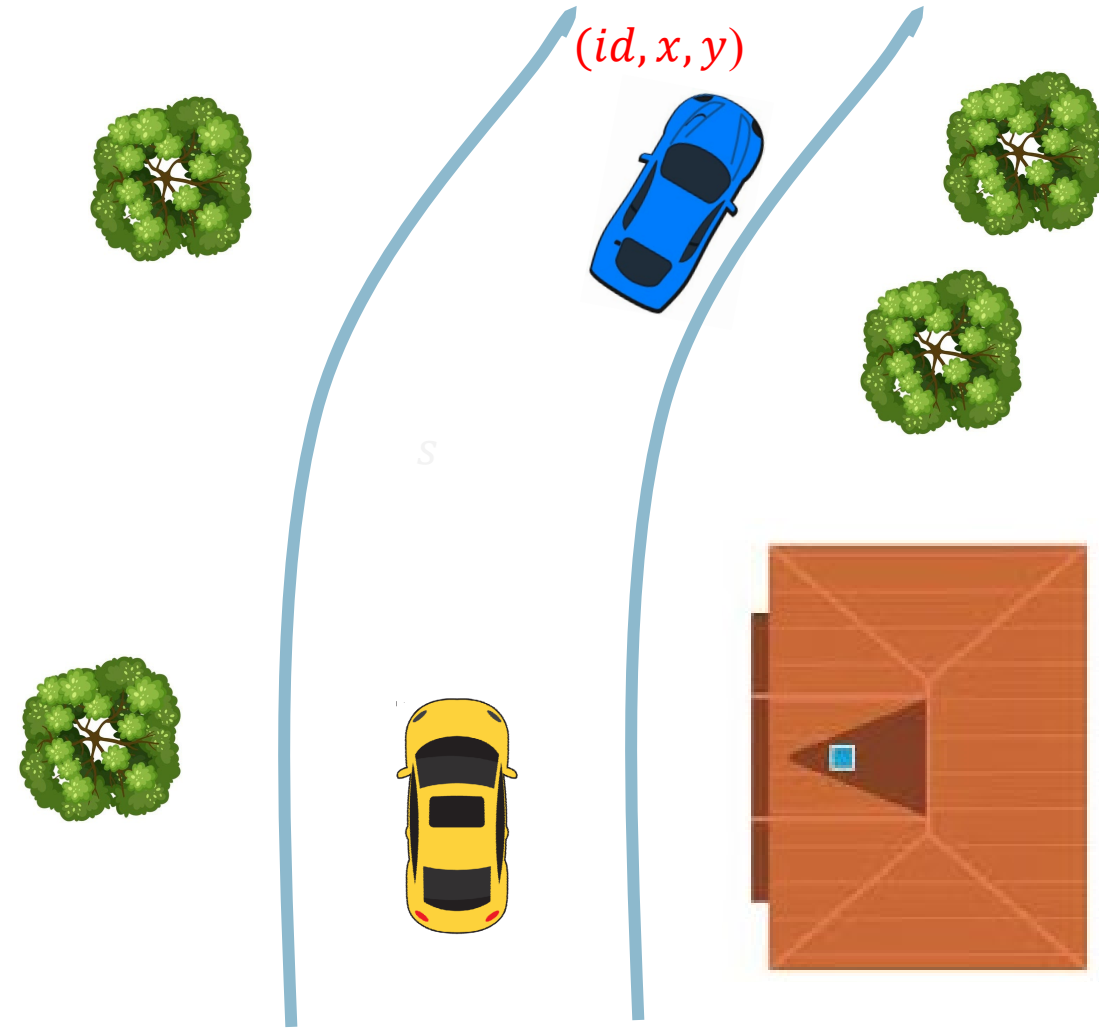
The object you are really interested in are the one you are going to crash with

Represent object in trajectory space

Sensors work in ego vehicle reference frame

RADAR (r, θ, v_x, v_y)

LiDAR/Camera (id, x, y)

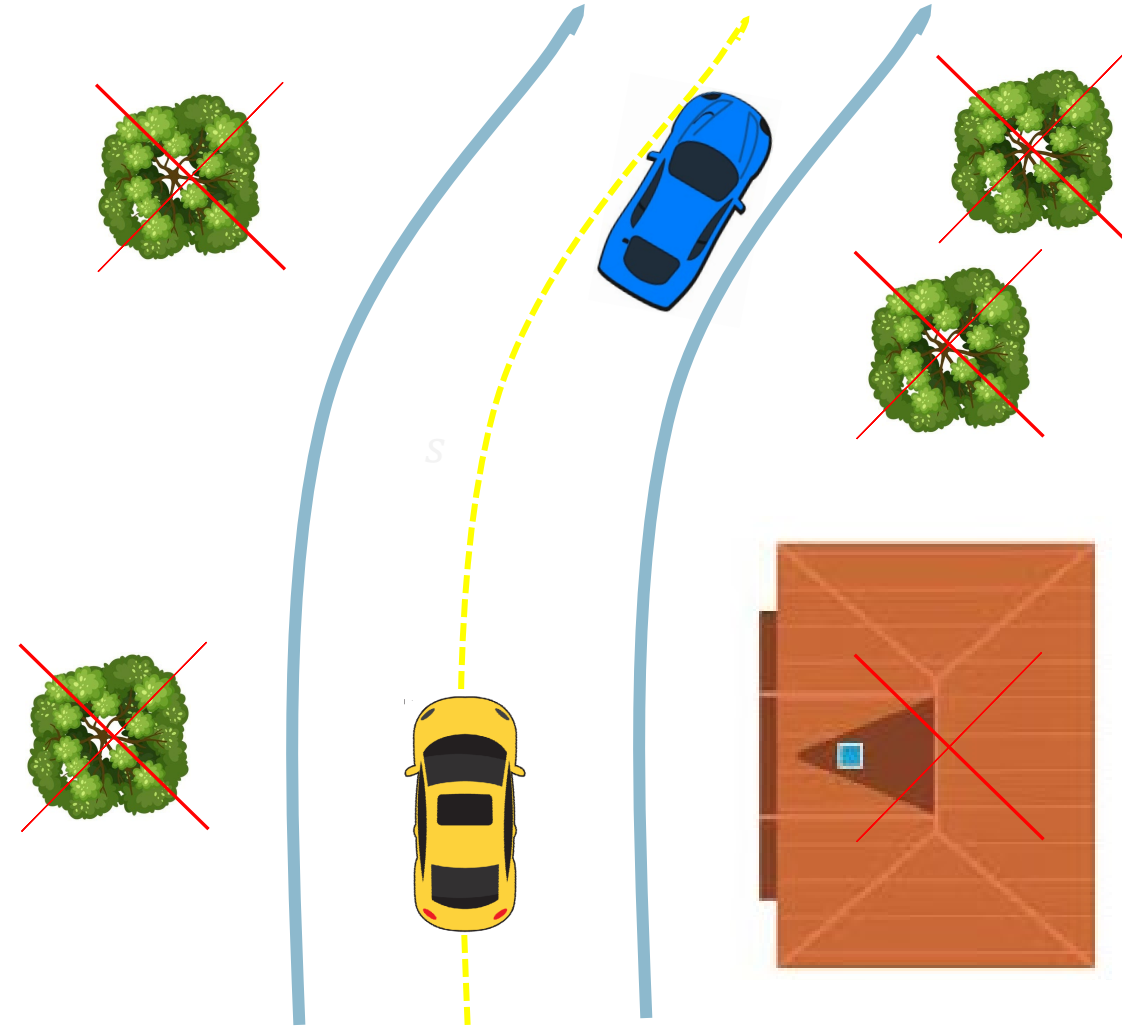


Seeing the others

On the road

Vision provide us the road shape

Everything outside of the boundaries is discarded

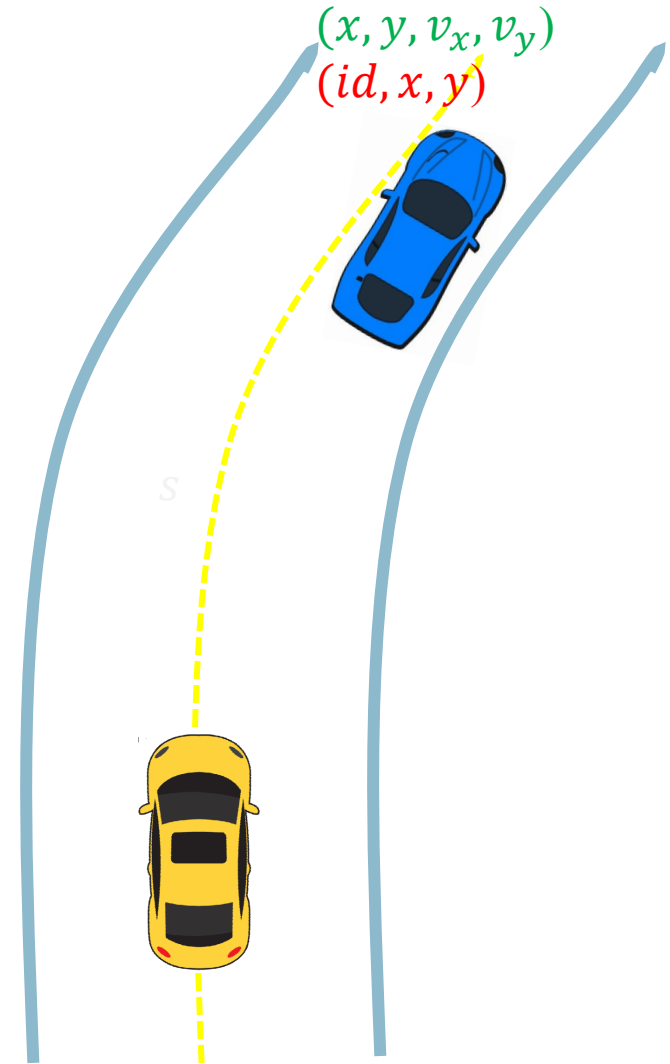


Seeing the others

Fusion

First we linearly combine the two sensors

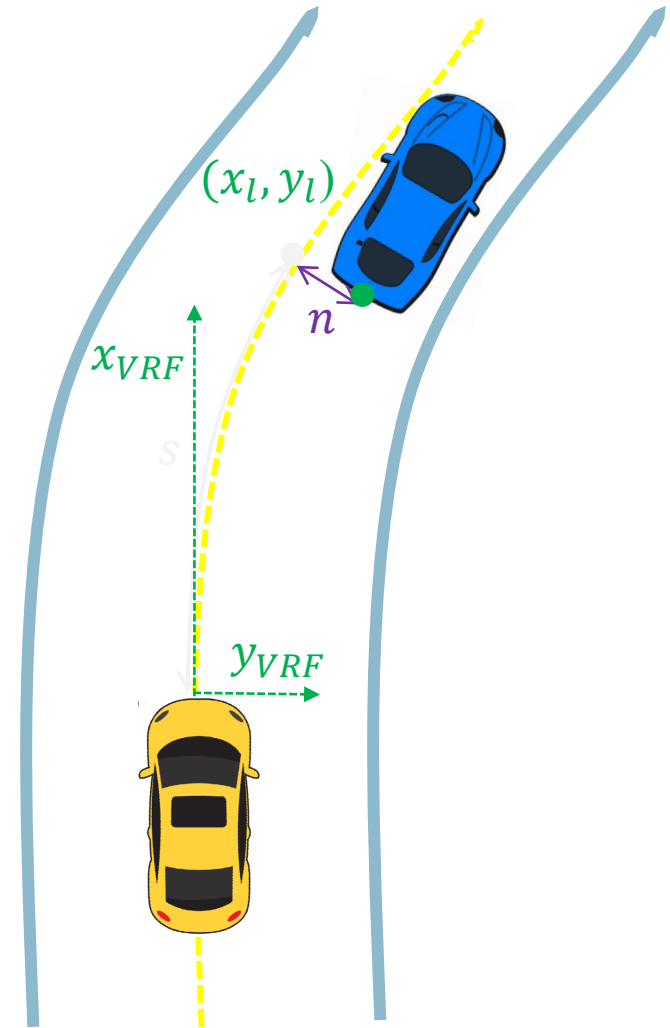
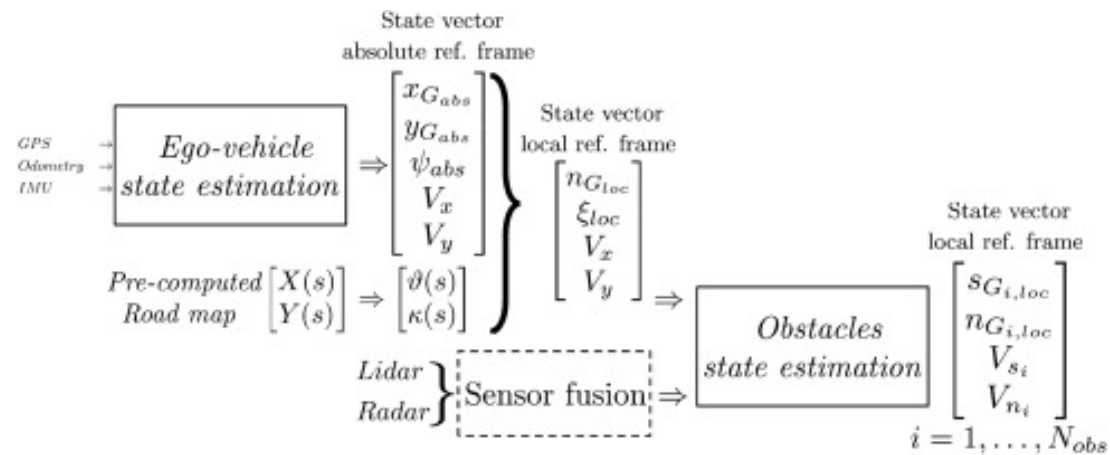
$$\begin{cases} x_{F_j}^{VRF} = 0.8 x_{l,o_i}^{VRF} + 0.2 x_{r,o_j}^{VRF} \\ y_{F_j}^{VRF} = 0.8 y_{l,o_i}^{VRF} + 0.2 y_{r,o_j}^{VRF} \\ V_{x,F_j}^{VRF} = V_{x,r_j}^{VRF} \\ V_{y,F_j}^{VRF} = V_{y,r_j}^{VRF} \end{cases}$$



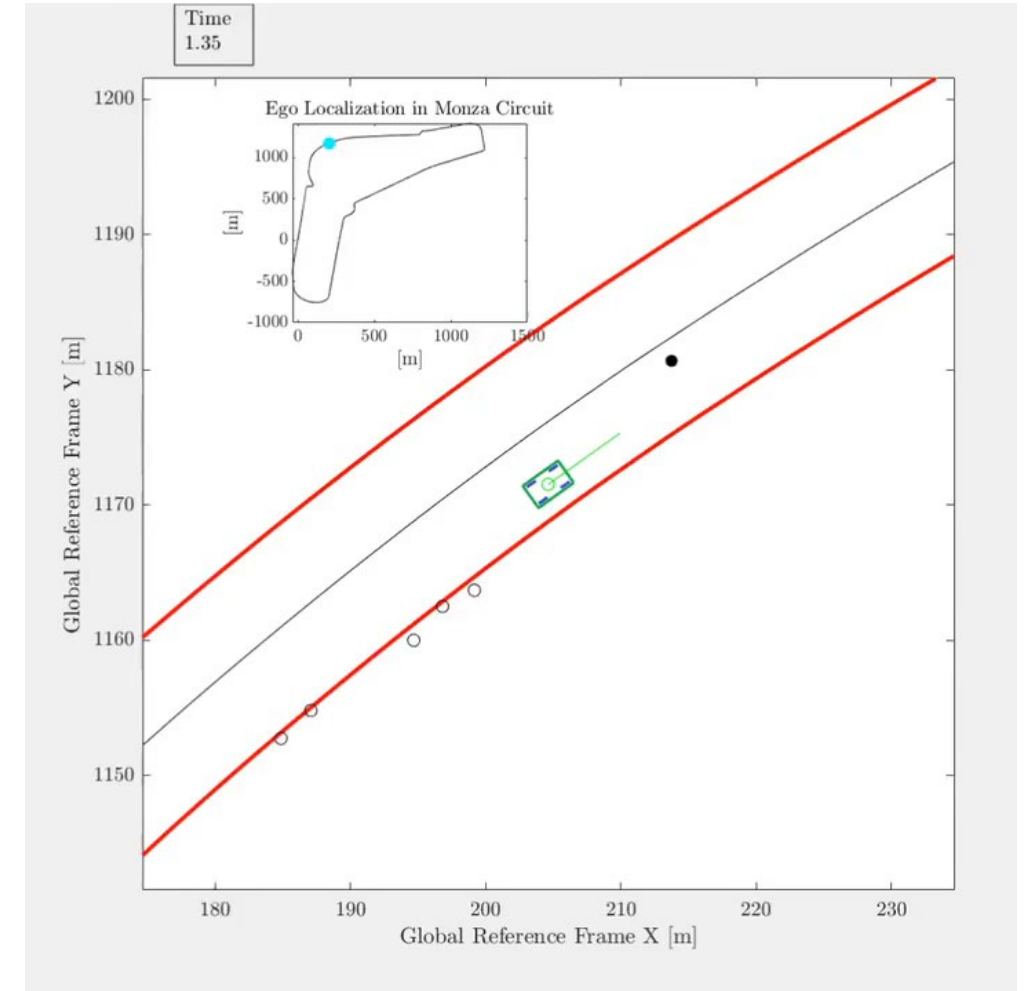
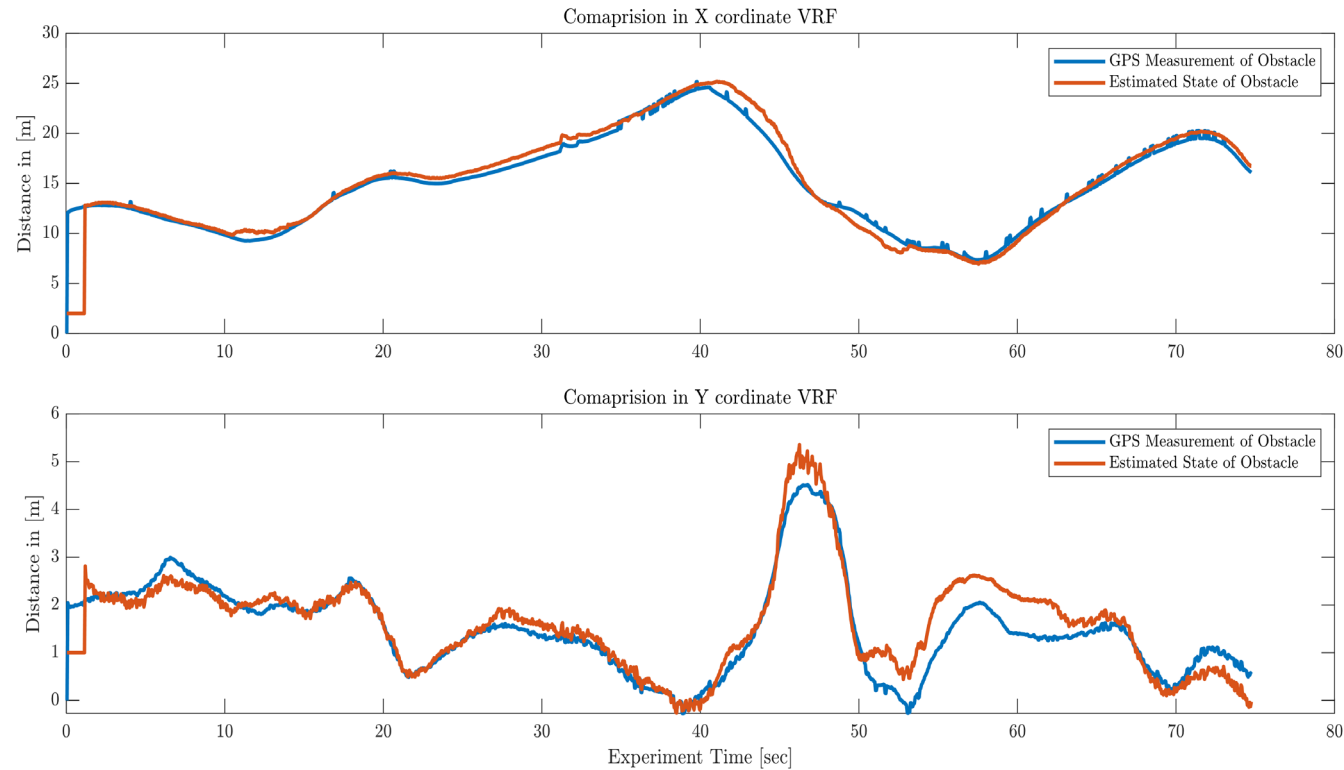
Seeing the others

Joint state

Then we represent the obstacle referenced to the ideal ego vehicle trajectory



Seeing the others



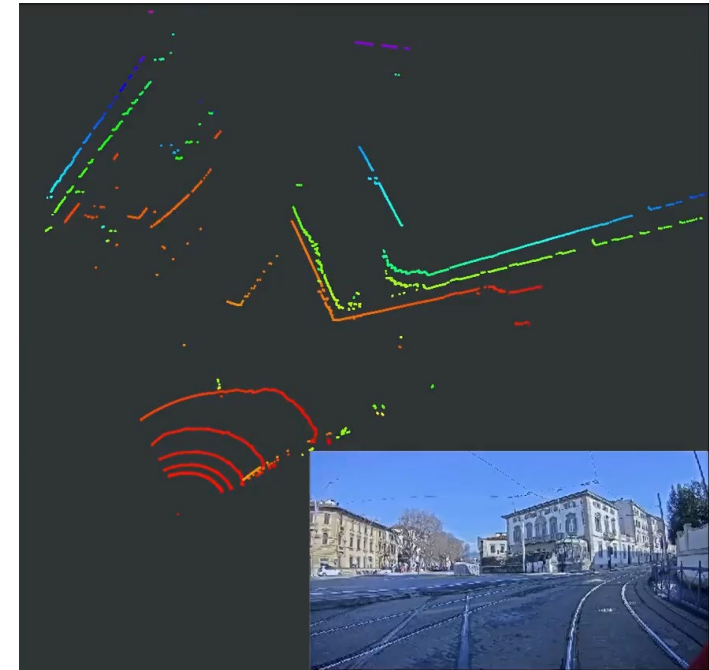
Seeing the others

LiDAR data

High-resolution LiDAR (128 planes)

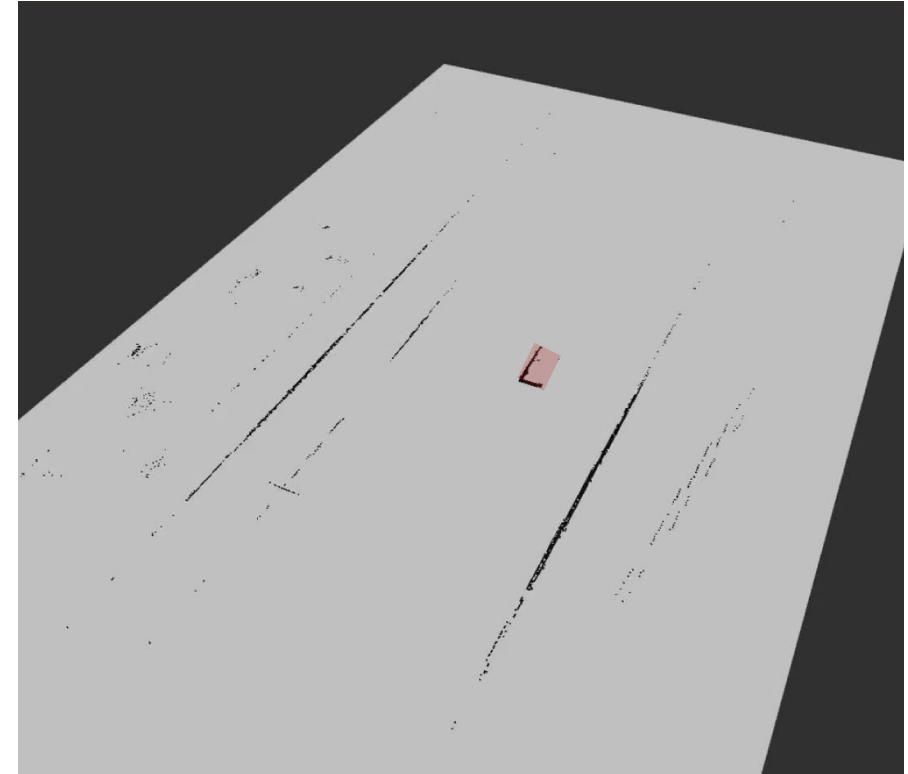
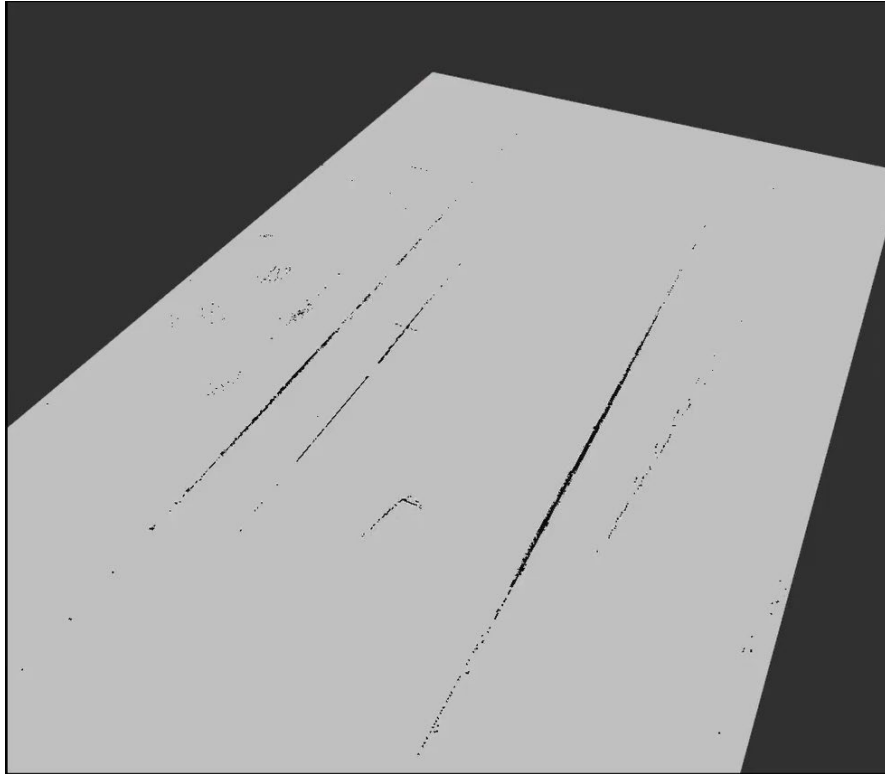


Low resolution LiDAR (8 planes)



Seeing the others

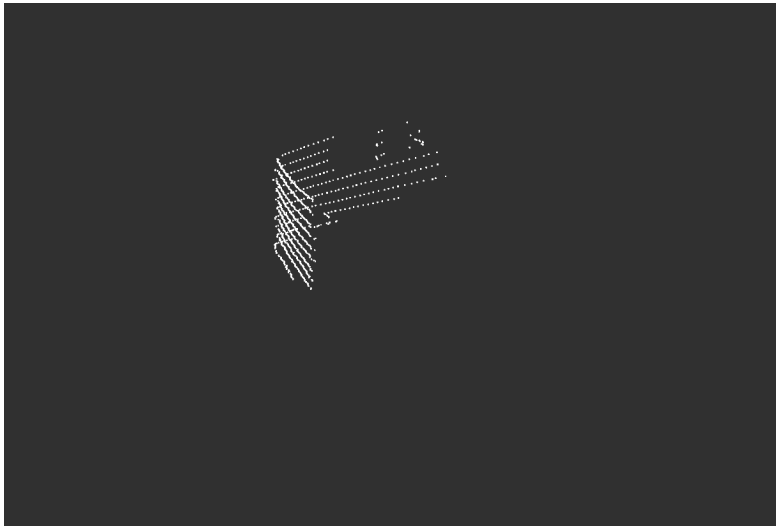
Occupancy Grid



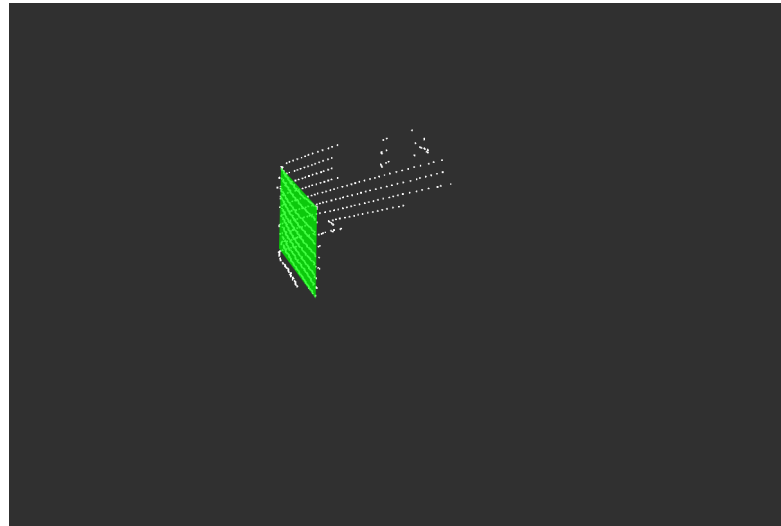
Seeing the others

3D Bounding Boxes

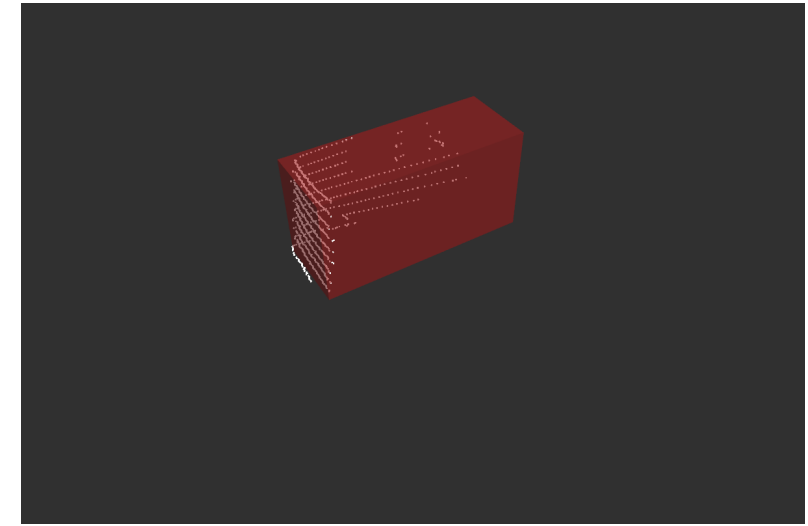
Cropped PointCloud



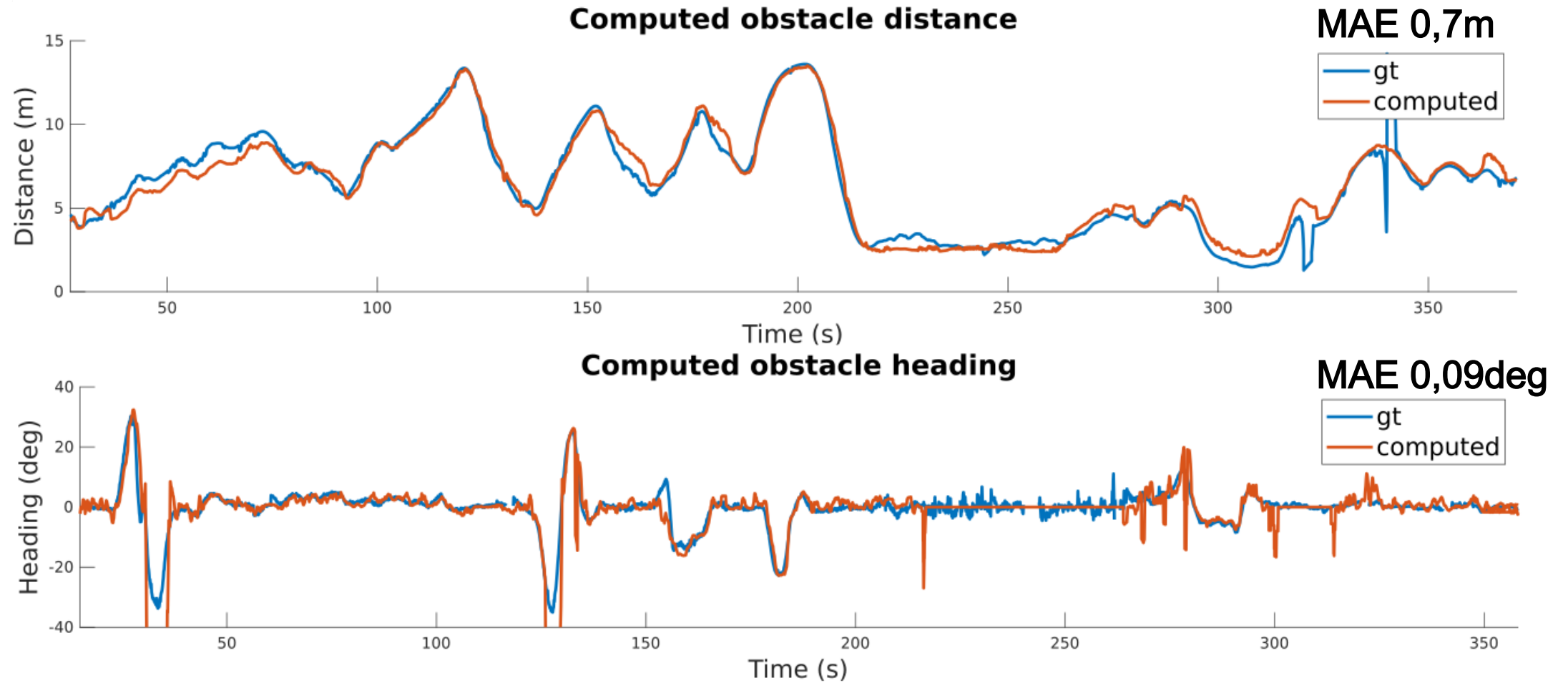
Plane Fitting



Bounding Box Fitting



Seeing the others



Cooperative Connected Vehicles Perception

02

Autonomous Vehicles limits

Expensive

Subject to
malfunctions

Fragile



Acceptability

Regulations

Explainability

Accountability

Instrumented Public Transportation

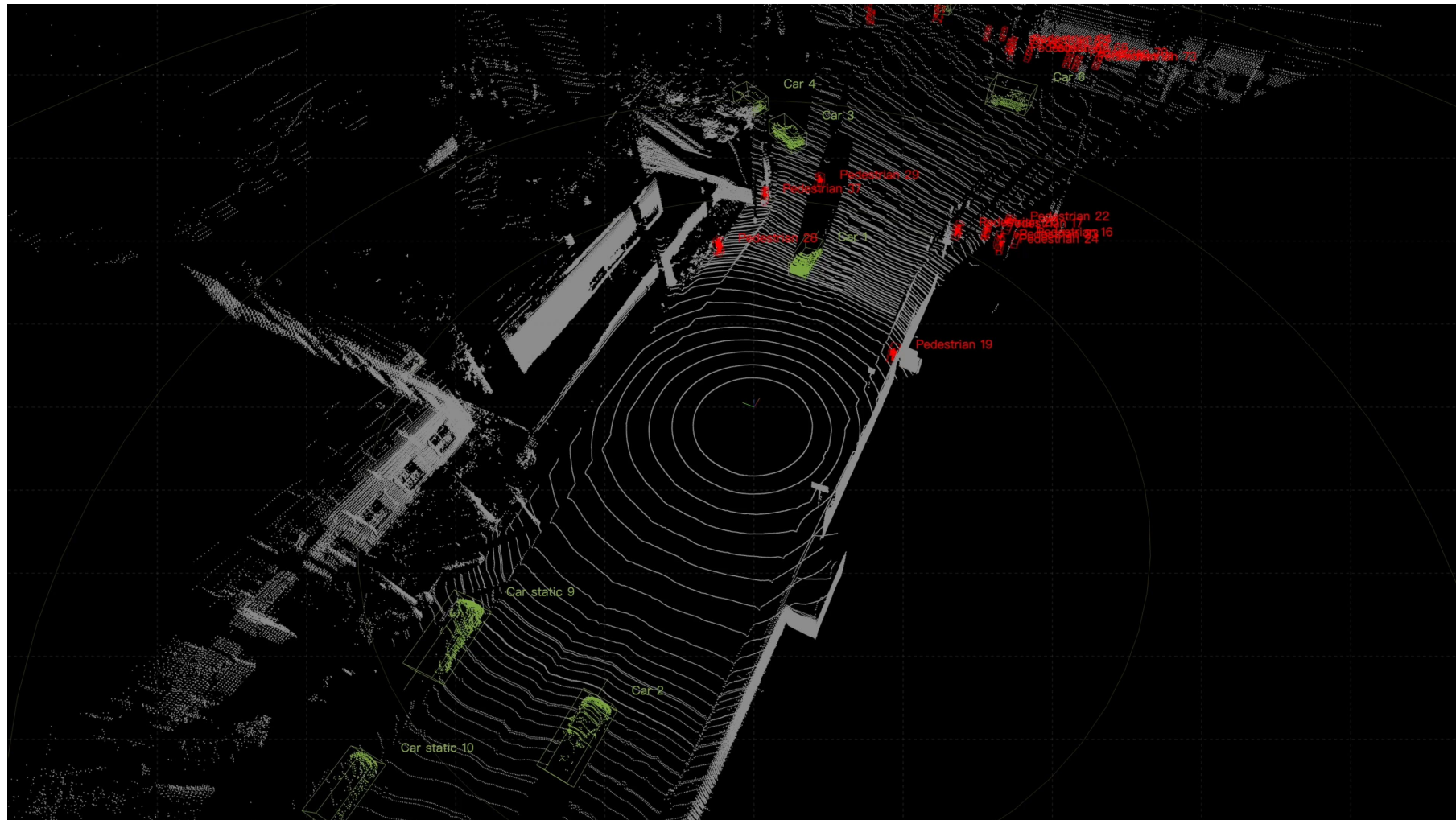
90/91 Milan



Florence Tram Line



From a higher point of view

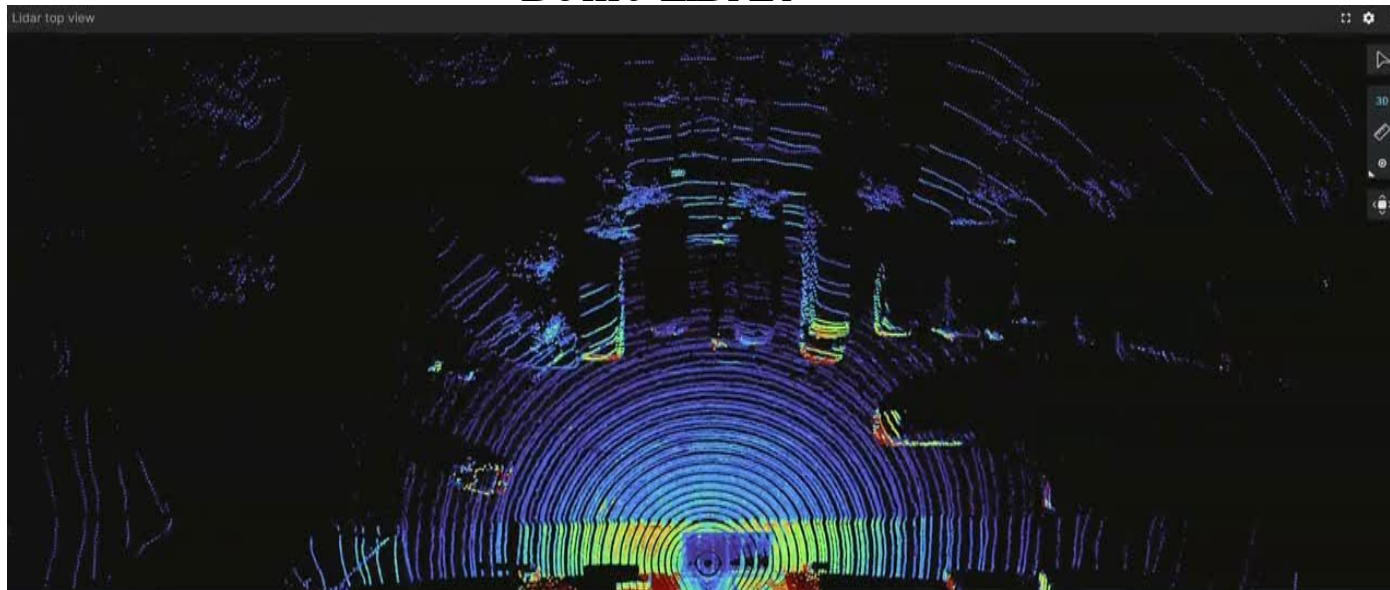


Xiang, H., Zheng, Z., Xia, X., Xu, R., Gao, L., Zhou, Z., Han, X., Ji, X., Li, M., Meng, Z. and Jin, L., 2024, September. *Real*: a large-scale dataset for vehicle -to-everything cooperative perception. In *European Conference on Computer Vision* (pp. 455-470). Cham: Springer Nature Switzerland.

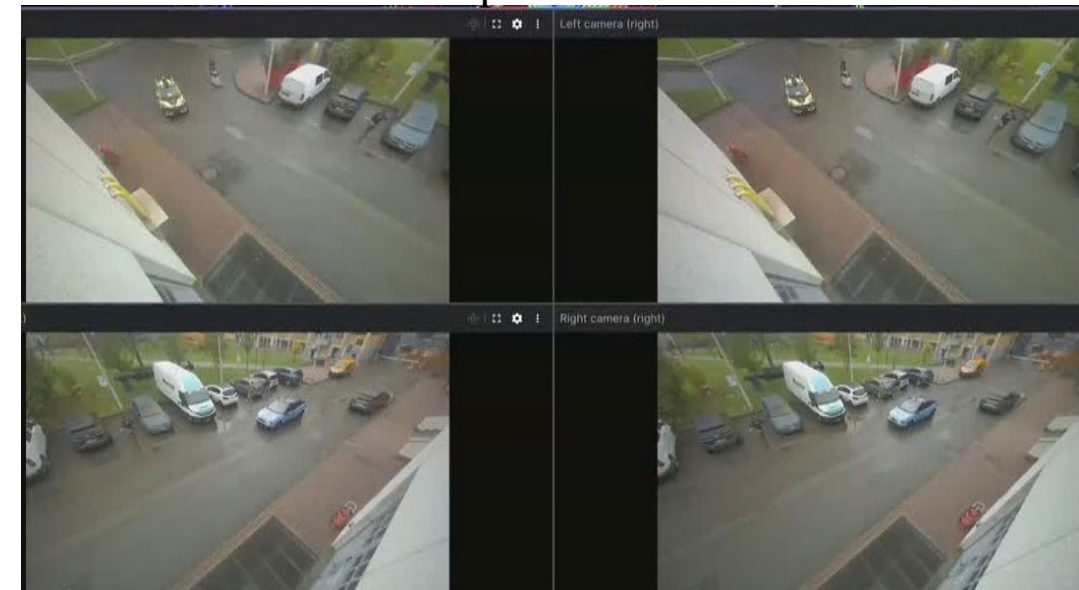
From a higher point of view

PoliMi Test Area

Dome LiDAR

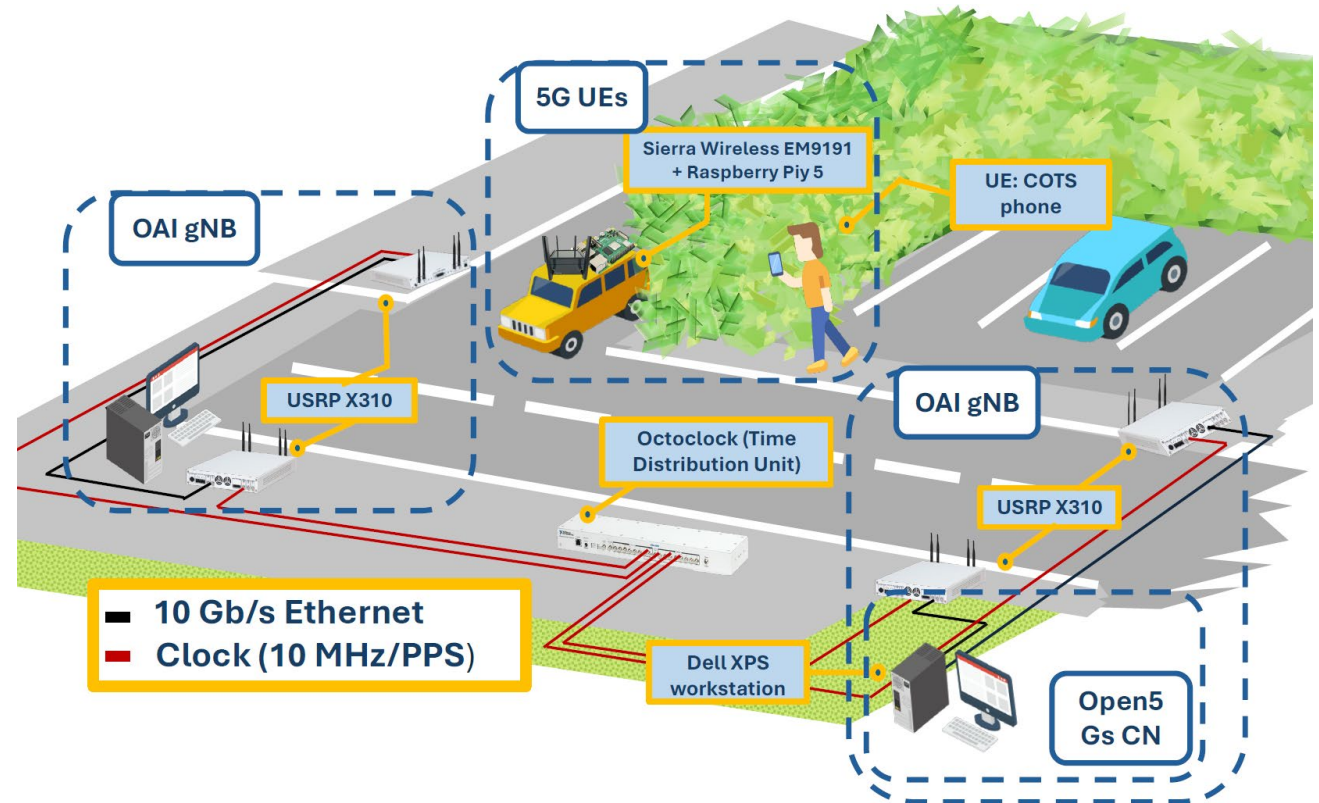
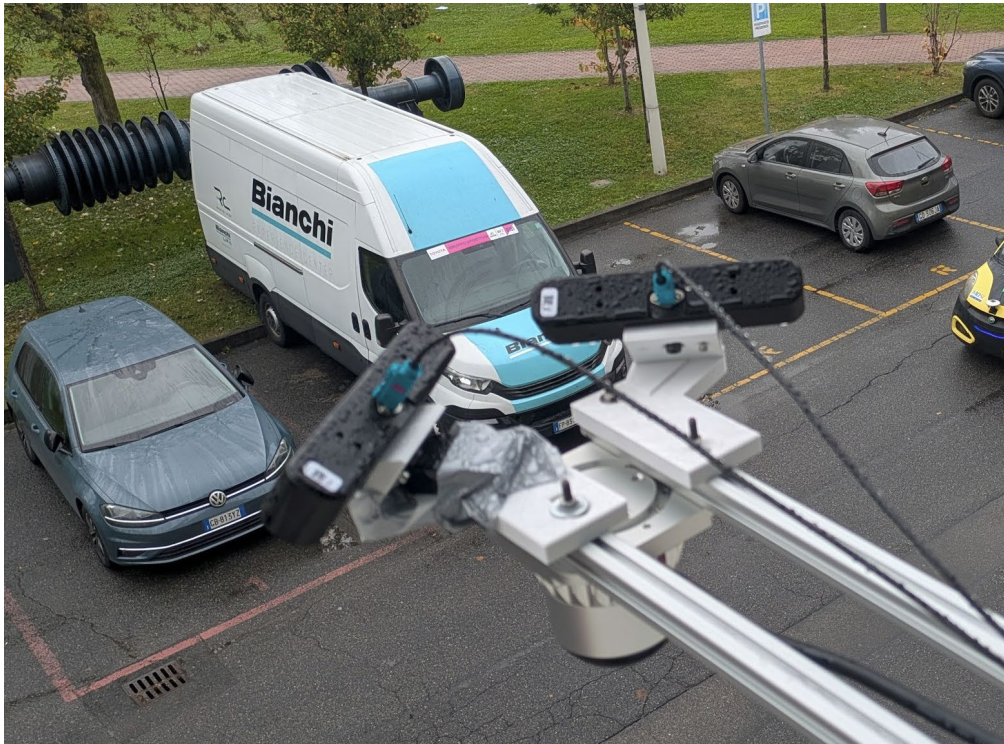


Multiple Cameras



From a higher point of view

PoliMi Test Area



Testing Scenarios

Vulnerable Road Users Protection



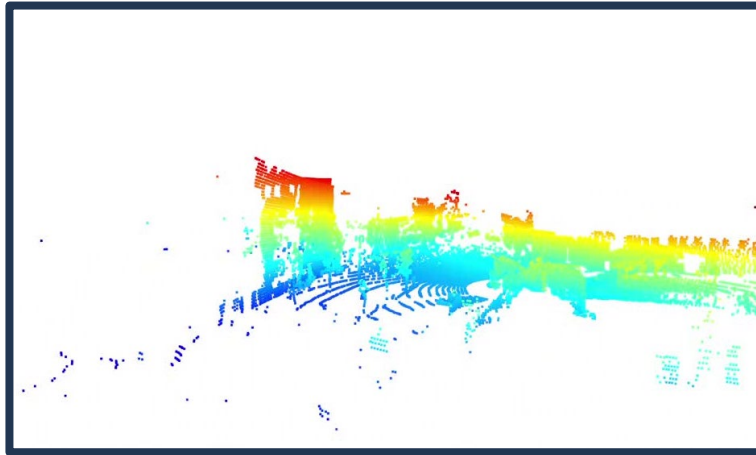
Autonomous parking



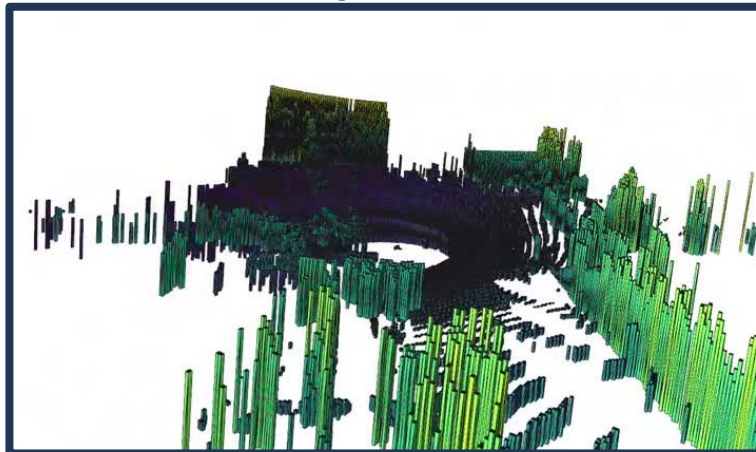
Limiting the power

2.5D grid -maps

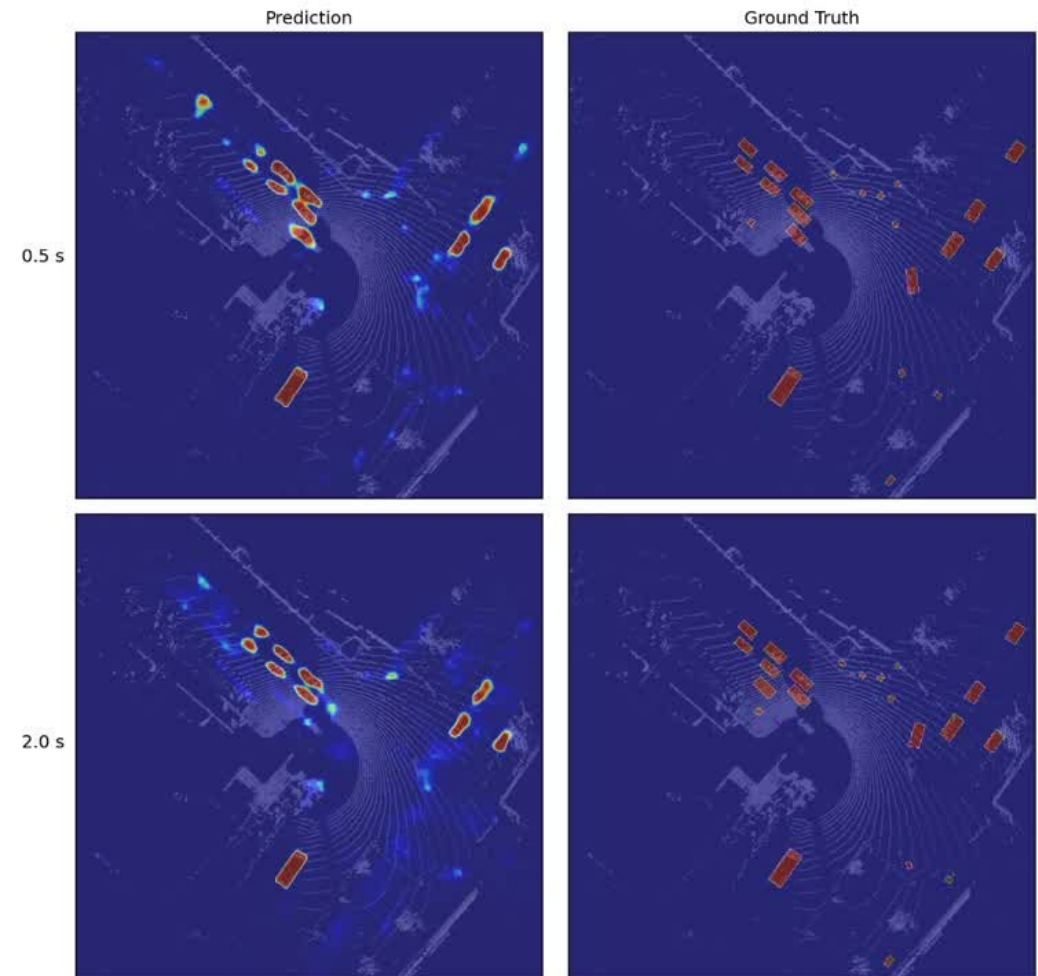
PointCloud



Height Map

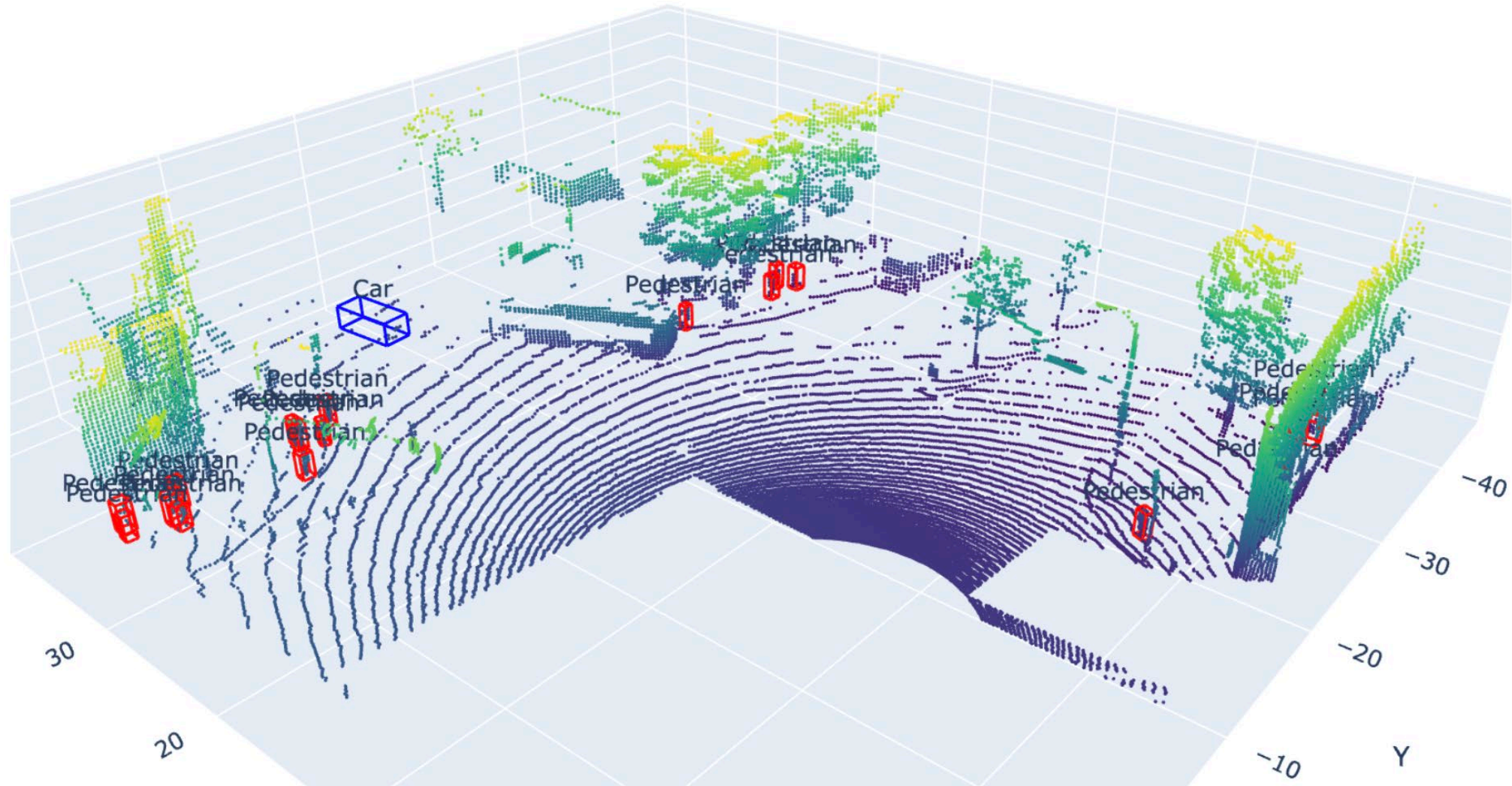


Detection and Prediction



Limiting the power

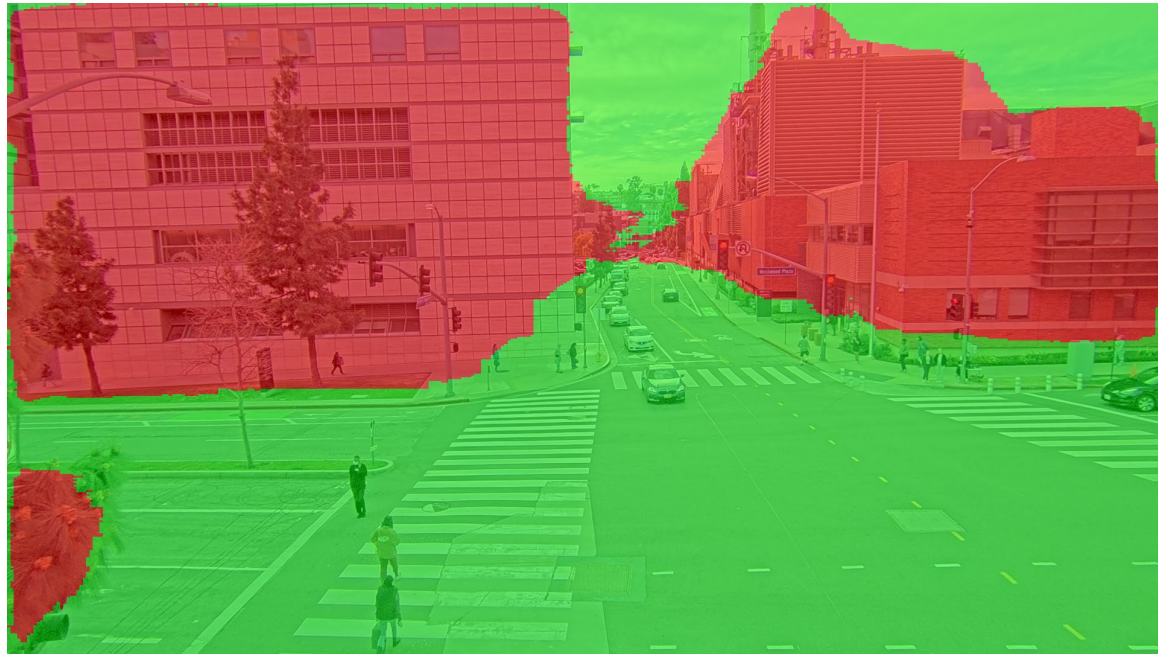
Context -aware filtering



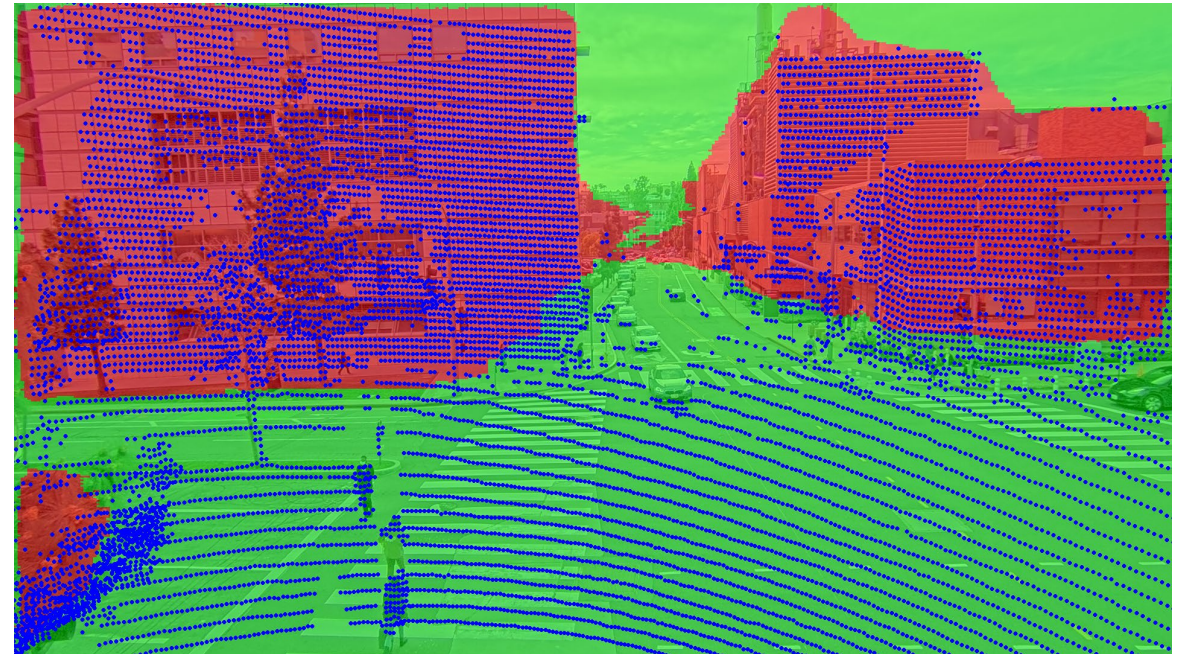
Limiting the power

Context -aware filtering

Static object not of interest



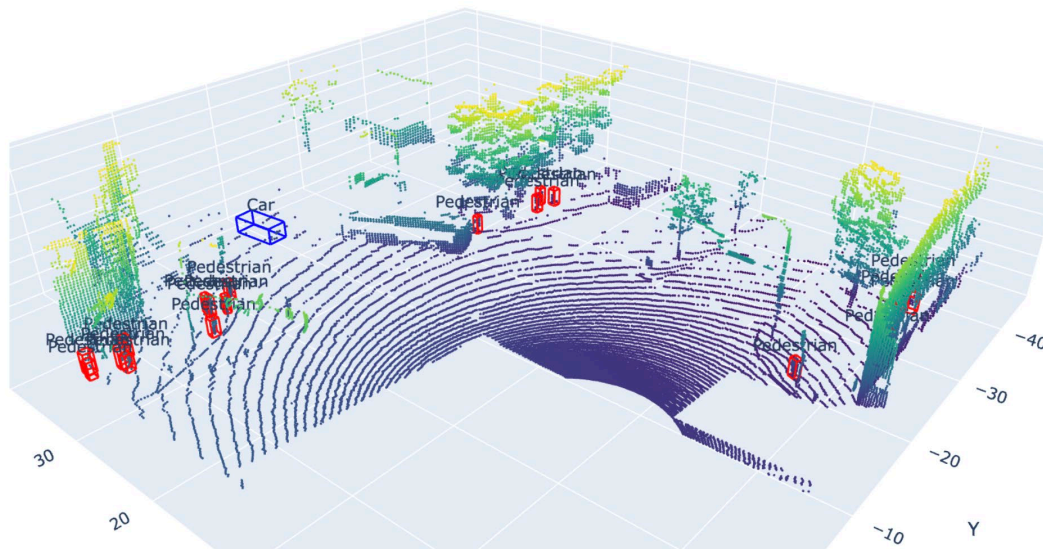
Filtered PointCloud



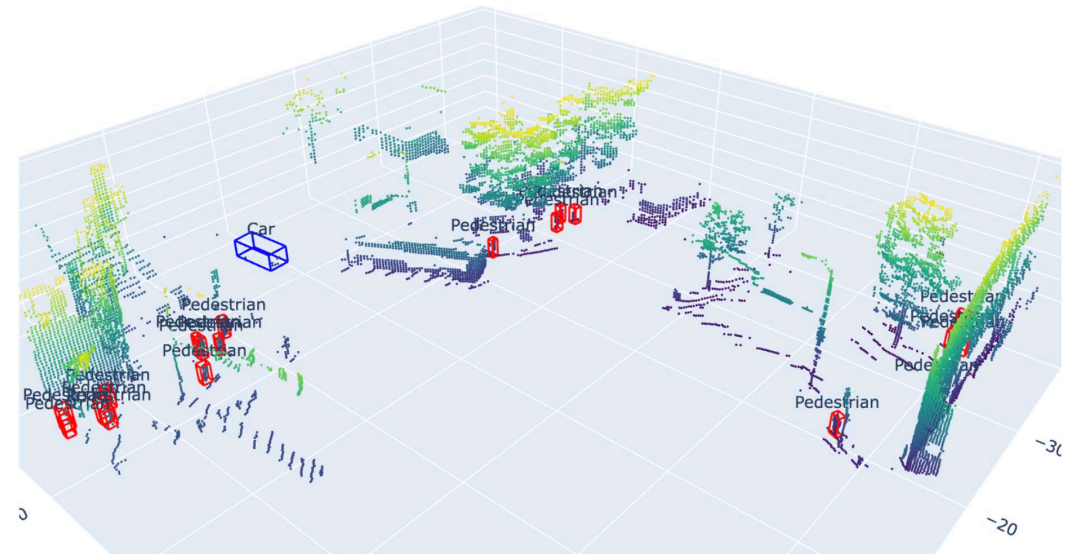
Limiting the power

Context -aware filtering

Original PointCloud



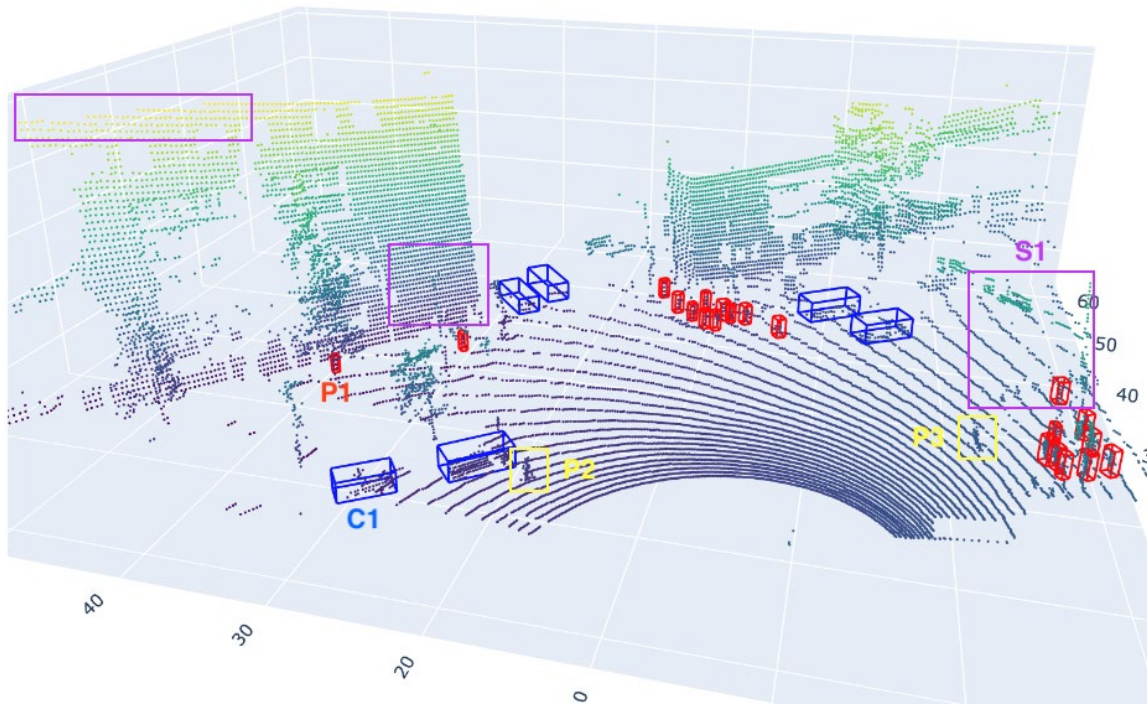
Ground removed



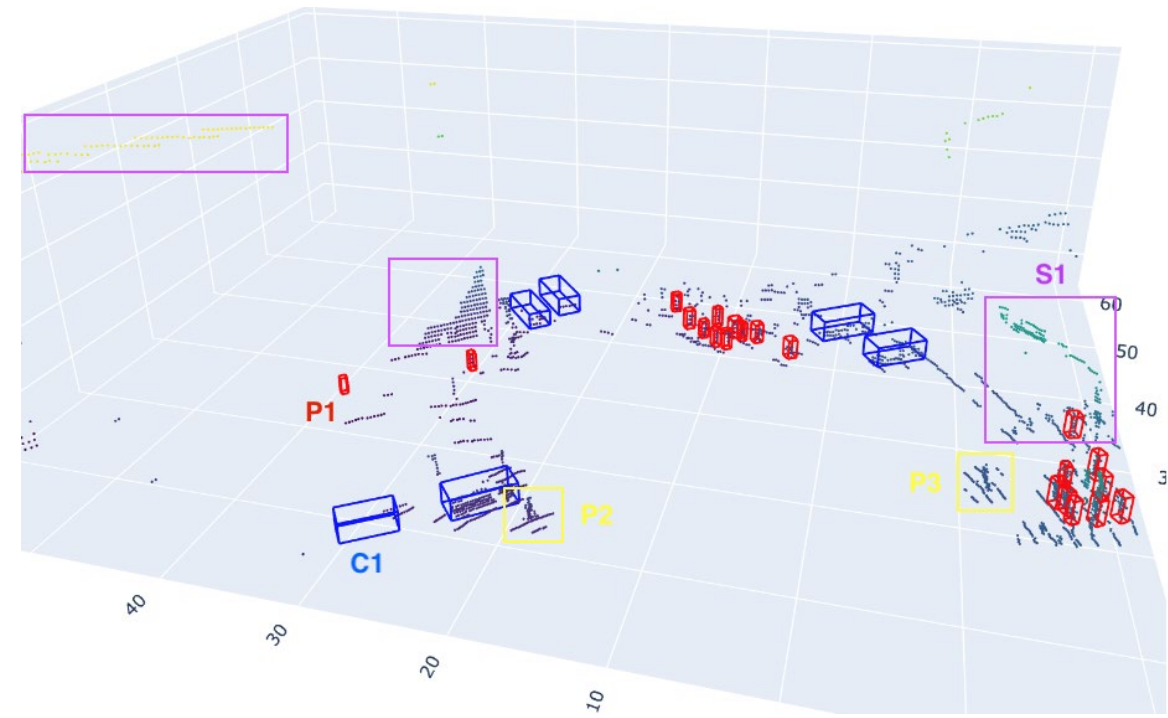
Limiting the power

Context -aware filtering

Original PointCloud



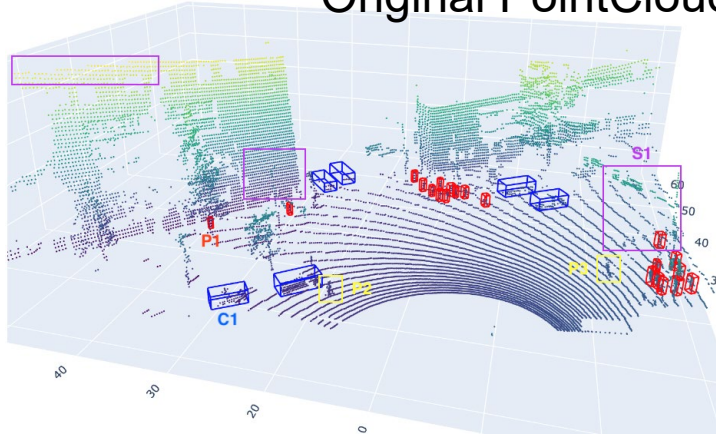
Filtered PointCloud



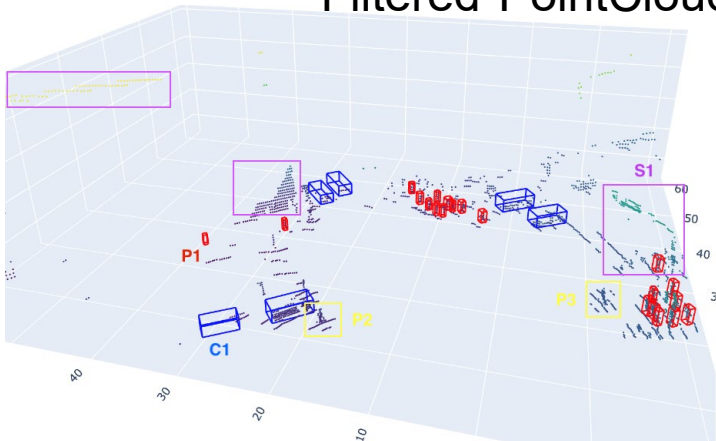
Limiting the power

Context -aware filtering

Original PointCloud



Filtered PointCloud



Inference Type	PC Size	Model	Car@0.50	Ped@0.35	mAP
No Filtered Inference	19294	MVXNet1	79.66	52.32	65.99
		MVXNet2	80.63	48.25	64.44
		SECOND	78.78	40.83	59.81
Static Object Removal Filter	13672 (-29.13%)	MVXNet1	75.86 (-3.8%)	50.03 (-2.3%)	62.95 (-3%)
		MVXNet2	76.60 (-4%)	44.77 (-3.5%)	60.69 (-3.7%)
		SECOND	74.83 (-3.9%)	36.72 (-4.1%)	55.77 (-4%)
Object-aware Ground Plane Filter	11011 (-42.93%)	MVXNet1	74.30 (-5.3%)	52.91 (+0.6%)	63.61 (-2.4%)
		MVXNet2	76.94 (-3.7%)	46.93 (-1.3%)	61.93 (-2.5%)
		SECOND	76.02 (-2.8%)	39.64 (-1.2%)	57.83 (-2%)
Combined Filters	5016 (-74.00%)	MVXNet1	70.69 (-8.9%)	50.63 (-1.7%)	60.66 (-5.3%)
		MVXNet2	72.79 (-7.8%)	43.27 (-5%)	58.03 (-6.4%)
		SECOND	72.48 (-6.3%)	35.56 (-5.3%)	54.02 (-5.8%)



Other research Areas

03

International Competitions

International
Field Robot Event

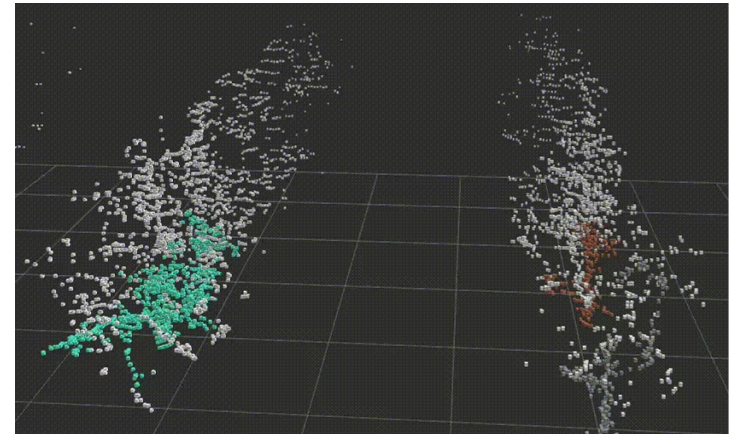


Agri - Robotics

Autonomous Field Coverage



Semantic Vineyards SLAM



Fruit Picking



