
Uni-DAS

15. Workshop
Fahrerassistenz
und
automatisiertes Fahren

FAS 2023



Copyright Uni-DAS e. V.

Alle Rechte, auch das des auszugsweisen Nachdrucks, der auszugsweisen oder vollständigen Wiedergabe, der Speicherung in Datenverarbeitungsanlagen und der Übersetzung, vorbehalten.

ISBN: 978-3-941543-74-4

Titelbild: Fahrzeuge UNICARagil, gefördert vom Bundesministerium für Bildung und Forschung

Uni-DAS e. V.

Otto-Verndt-Straße 2

64287 Darmstadt

Uni-DAS

15. Workshop Fahrerassistenz und automatisiertes Fahren

FAS 2023

24. – 26.10.2023

Kloster Bonlanden, Berkheim

Vorwort

In diesem Jahr findet der 15. Uni-DAS e.V. Workshop „Fahrerassistenz und automatisiertes Fahren“ statt. Es gibt also etwas zu feiern! Wie dem runden Geburtstag des Workshops zu entnehmen ist, hat er sich mittlerweile als fester Bestandteil der „FAS-Community“ im deutschsprachigen Raum etabliert und bietet Expertinnen und Experten aus Forschung und Wirtschaft ein einzigartiges Forum zur interdisziplinären Diskussion. Der Workshop-Charakter wird durch Verzicht auf Parallelveranstaltungen, Zeit für Diskussionen innerhalb und außerhalb der wissenschaftlichen Sessions, Kleingruppenarbeit zu einem aktuellen Thema, sowie die Beschränkung der Zahl der Teilnehmenden geprägt. Dadurch ist der Workshop etwas Besonderes, was sich nicht zuletzt in der fortwährenden Überbuchung der Veranstaltung widerspiegelt.

Zur Zeit des ersten Workshops im Jahr 2002 waren Fahrerassistenzsysteme erst kurz im Markt, und die wenigen verfügbaren Komfortfunktionen waren Oberklassefahrzeugen vorbehalten. Inzwischen sind diese Systeme in nahezu allen Fahrzeugklassen weit verbreitet. Heute sprechen wir über Level 3 Funktionen in Serienfahrzeugen. Wie geht es nun weiter? Wie kommen Systeme auf den Markt? Welche Geschäftsmodelle könnten erfolgreich sein? Über diese Fragen möchten wir mit den Teilnehmenden im Workshop diskutieren. Wir freuen uns auf Erfolgsmodelle, so dass wir auch auf unseren Straßen bald vermehrt automatisierte Fahrzeuge fahren sehen. Die einzelnen Vorträge des Workshops und Beiträge dieses Tagungsbandes beleuchten ein breites Spektrum des aktuellen Wissensstandes im Bereich der Fahrerunterstützung und des automatischen Fahrens.

Unser Dank gilt allen Helfern und Helferinnen bei der Organisation des Workshops sowie den Mitgliedern von Uni-DAS e.V.. Wir wünschen allen Teilnehmenden neue Einsichten, spannende Diskussionen und das Knüpfen neuer und Vertiefen alter Kontakte.

Berlin und Aachen

Meike Jipp und Lutz Eckstein

Inhaltsverzeichnis

Daten und Wahrnehmung

- | | | |
|---|---|----|
| 1 | From Vehicle Setup to Dataset Generation: A Holistic Approach to Long-Range Automated Valet Parking Development | 1 |
| | Arab, A. B. (Expleo) et al. | |
| 2 | Map Learning: Ein Skalierbarer Ansatz zur Automatisierten Erstellung von Trainingsdaten unter Verwendung von HD Karten und Mehrfachbefahrungen | 17 |
| | Bieder, F. (FZI) et al. | |
| 3 | Random-Finite-Set-basiertes Multisensor-Multiobjekttracking für automatisierte und vernetzte Fahrzeuge | 27 |
| | Herrmann, M. (U Ulm) et al. | |

Methoden der künstlichen Intelligenz

- | | | |
|---|---|----|
| 4 | Physics-informed Reinforcement Learning for Automated Merging in Dense Traffic | 43 |
| | Fischer, J. (KIT) et al. | |
| 5 | Reducing Ghost Detections Through Uncertainty Modelling for Automated Driving | 53 |
| | Hammam, A. (Stellantis) et al. | |
| 6 | Transfer Learning Techniques Using Simulation Data For Machine Learning Automotive Radar Systems | 67 |
| | Rutz, F. (TUM) et al. | |

Verhalten und Interaktion

- | | | |
|----|--|-----|
| 7 | Spieltheoretischer prädiktiver Regler für Interaktion im gemischten Verkehr | 77 |
| | Bouzidi, M.-K. (Continental AG) et al. | |
| 8 | Automatisiertes Kreuzungsmanagement im urbanen Mischverkehr mit Reinforcement Learning | 85 |
| | Klimke, M. (Robert Bosch GmbH) et al. | |
| 9 | Semantische Normverhaltensanalyse zur durchgängigen, formalen Verhaltens-spezifikation automatisierter Straßenfahrzeuge | 95 |
| | Salem, N. F. (TU Braunschweig) et al. | |
| 10 | Bedingte Verhaltensprädiktion interagierender Agenten | 111 |
| | Wirth, F. (KIT) et al. | |

Risiko und Sicherheit

- | | | |
|----|--|-----|
| 11 | Systematic Derivation of Use Case Clusters for a Generalized Low-Speed Automated Driving Function | 121 |
| | Berghöfer, M. (TU Darmstadt) et al. | |
| 12 | Derivation of quantitative risk acceptance criteria for automated driving systems | 139 |
| | Stellet, J. (Robert Bosch GmbH) et al. | |
| 13 | SURE-Val: Safe Urban Relevance Extension and Validation | 153 |
| | Storms, K. (TU Darmstadt) et al. | |
| 14 | Operator Monitoring zur Absicherung von Teleoperation und automatisiertem Fahren | 169 |
| | Herzberger, N.(RWTH Aachen) et al. | |

From Vehicle Setup to Dataset Generation: A Holistic Approach to Long-Range Automated Valet Parking Development

Amine Ben Arab*, Jing Gu[†], Hassan Mohammadi[‡], Mia Book[‡],
Thomas Brandmeier[‡], Alireza Ferdowsizadeh Naeeni[†]

Abstract: Sensor fusion is one of the trend topics in the automotive field, which aims to achieve robustness by relying on the information from different sensors. But to achieve further progress in this field, the availability or generation of multimodal data under different contexts is a vary important topic. In this paper we focus in the description of the necessary methodologies to generate such data considering Long-Range Automated Valet Driving scenarios, involving a combination of urban and indoor driving scenarios, we devised two distinct methods. For the initial approach, we utilized the installations present in the ISAFE Indoor Testing Facility at CARISSMA (Technische Hochschule Ingolstadt) to gather data for situations in severe environments and supply the necessary references. The second approach involves producing data from real-life situations using the prototyping vehicle of Expleo Germany GmbH. To accomplish this, we equipped the vehicle with various sensors and examined techniques from the latest technological developments to decrease the burden of dataset generation while meeting sensor fusion’s demands. In this work, we present both approaches and our results for the methods employed to establish our data generation pipeline.

Schlüsselwörter: AVP, Sensors, Sensor-2-Sensor Calibration, Time Synchronization, Reference systems.

1 Introduction

Long-range Automated Valet Parking (LAVP)[1] is an extension of the AVP function. This service not only drives and parks autonomously in a parking facility like typical AVP but also considers driving the vehicle from a drop-off zone far from the parking facility’s entry. Therefore, this function encompasses more diverse scenarios. On one hand, the system must cope with the challenges of urban environments, including dynamic conditions and various weather elements such as rain and fog. On the other hand, it must also manage the intricate maneuvering involved in enclosed parking situations.

Throughout the years, multiple datasets [2...13] have been made available to the public for AD use cases due to the significant time and cost required to generate a dataset.

*Expleo Germany GmbH, Wilhelm-Wagenfeld-Str.1-3, 80807 München (e-mail: amine.ben-arab@expleogroup.de).

[†]Expleo Germany GmbH, Salufer 8, 10587 Berlin (e-mail: name.lastname@expleogroup.com).

[‡]Technische Hochschule Ingolstadt, CARISSMA ISAFE, Esplanade 10, Ingolstadt (e-mail: name.nachname@carissma.eu).

These datasets have provided increasing sensor [2] and scenario [3] diversity, enabling the research community to rapidly prototype for scene understanding. However, to the best of our knowledge, there is currently no open dataset that accurately reflects the various scenarios that describe our particular use case.

In this paper, we present two approaches to data acquisition that focus on a diverse data set that can be collected in a time-efficient manner without extensive human intervention. The initial approach is achieved by capturing mock scenarios within a controllable indoor environment. By utilizing an indoor positioning system (IPS) to determine the location of individual objects, the labeling procedure can be partially automated, rendering it dependable even during undesirable weather conditions. The second approach involves gathering data in real-world environments. In this regard, we present the sensor setup we implemented in our prototyping vehicle, along with our sensor calibration and time synchronization methods. We also introduce our reference system for the vehicle's ego-localization task indoors.

2 Related Works

Over the last years, dataset generation has been a focus of research in AD, especially multi-modal datasets have become more and more important [3]. This is due to the rapid progress in sensor development in the LiDAR and radar fields. There are numerous techniques for creating a dataset to train and test driving-related algorithms, including collecting data in real-world settings, simulated environments, or performing mock scenarios. Real-world data acquisition provides realistic and authentic sensor data with accompanying noise but poses the risk of generating edge cases or critical scenarios. Additionally, annotating the data is a time-consuming and potentially incomplete process. For example, under rainy conditions, Lidar can't model the whole object of interest and therefore, it is challenging to generate a 3D bounding Box that limits this object. With simulated data, it is easy to generate critical and rare scenarios as well as annotations. However, the sensor data obtained from LiDAR or radar might not be comparable with sensors used in the actual vehicle. The Table 1 summarizes various real-world multi-modal datasets as well as datasets that include adverse weather conditions. It is evident that the majority of the datasets were collected during favorable conditions (cloudy or clear) or did not label all collected adverse conditions [4]. Only a few, such as Rain WCity, focus on collecting data during adverse weather conditions [5]. Additionally, parking scenarios are either not explicitly mentioned in the literature or are absent from most datasets. Different types of sensors are being researched to overcome the challenges of AD. However, to meet the automotive industry's standards, factors such as reliability, durability, and cost-effective manufacturing and maintenance must be considered. In this context, the sensors with the potential to solve high-level AD, found in the datasets listed in Table 1, are radar, camera, and LiDAR. These different sensors can be used as a redundancy solution, to enhance the reliability of the AD system, as well as a complementary solution to extend the continuity of the AD system.

To fuse the information from the sensors properly, a calibration is a necessary step. To ensure spatial consistency between sensor data, it is required to transform each sensor's readings into a common coordinate system, known as extrinsic calibration. It can be accomplished through different methodologies. For example, the sensor data can then

be optimized through manual alignment, as seen in Radiate [4] and partially in Nusences [2] and EU long term [6], or the transformation between sensors can be estimated automatically using data from the sensors, which consists in target-based, and targetless sensor-2-sensor calibration [7]. Target-based calibration utilizes a specially designed target with known dimensions to extract features from each sensor. The chessboard target is frequently used to extract its corners as a 3D pattern. in Pixset [8], To estimate the transformation matrices between the cameras, the authors employed a closed-loop optimization method and used the perspective-n-point method for calibration between the camera and LiDAR. For chessboard-based calibration on the Kitti dataset, they utilized their later work [9]. Additionally, a unique target was created for the calibration of LiDAR, camera, and Radar at Tu-Delft [10]. A further investigation of the calibration method is presented in section 3.2. Targetless calibration consists of two sub-methods: Motion-based and feature-based calibration. Motion-based [11] calibration estimates the vehicle's trajectory from each sensor during various driving maneuvers and matches those trajectories to estimate the transformation between the sensors. This technique is commonly used for proprioceptive sensor calibration, including IMU or GNS. In the KITTI dataset, the hand-eye method was applied to calibrate IMU and LiDAR [9]. The feature-based calibration [12] extracts similar features from each sensor within the shared scene and estimates the transformation matrix by identifying correspondences between these features. In the A2D2 dataset [13], calibration between the camera and LiDAR is enhanced by utilizing edge correspondences.

Time synchronization of sensors is critical to the sensor fusion process. It is essential to maintain temporal consistency between sensor data. Simultaneous sensor triggering is the most accurate method for fusing sensors by simultaneously registering events on each sensor. In the KITTI dataset, cameras are triggered simultaneously with LiDAR by a reed contact in the LiDAR that activates the camera when the mechanical LiDAR's scanner faces forward. However, simultaneous sensor triggering is not feasible for the various automotive-grade sensors. The alternative is to synchronize the sensors to a common clock, and thus this method maintains the temporal correlation of events between sensors. this method can be achieved either at the software level or at the hardware level. In the EU long-term dataset, the LiDARs are synchronized at the hardware level by synchronizing the clocks of each LiDAR with the acquired GPS signal. Another method is to timestamp the sensor data with the time it arrives at the data logger. This method may be less accurate than others as it includes a non-deterministic delay in data transmission.

As most of the datasets are generated outside, an RTK GPS is usually used for generating reference positioning of the vehicle. It's an aviation-grade positioning system and is considered the most accurate positioning system for outdoor use cases.

Data Annotation consists of labeling the different static and dynamic objects in the vehicle's surroundings, which are relevant to the AD function. It is done manually with the help of annotation tools [14][9] [3]. It's a rigorous task and hence very time-consuming as the precision of the labels has a big impact on the targeted task. In Cityscape [14], the semantic labels generation for each frame required 1.5 hours on average. however, there are some efforts to generate these labels automatically and reduce the labeling burden. Human supervision is still needed to validate or correct the proposed labels. For example, in the aiMotive dataset [15], a method was used to search for possible candidates in a sequence of data. the candidate's labels are further enhanced with an object tracking

algorithm, where they are optimized recursively.

Table 1: State-of-the-art datasets. Abbreviations scenarios: (U)urban, (SU)suburban, (H)highway and (P)parking; conditions: (G)good, (N)night, (R)rain, (F)fog and (S)snow; Sensors: (M)mono/(S)stereo (C)camera, (M)mechanical/(SS)solid-state (L)LiDAR and (R)radar; synchronization (Sync.): (H)hardware and (S)software; calibration: (M)manual, (TL)target-based and (TB)target-less; (-)no information available or not available in the dataset.

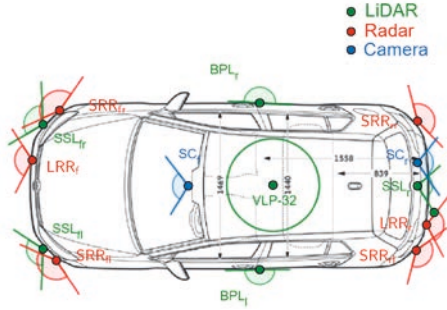
Dataset	Scenarios				Conditions				Sensors			GPS/IMU		Annotation			Classes	Sync.	Calib.
	U	SU	H	P	G	N	R	F	S	C	L	R		2D	3D	Sem.			
KITTI [9]	✓	✓	✓	-	✓	-	-	-	-	1 SC	1 ML	-	✓	-	✓	✓	28	H/S	TB
nuScene [2]	✓	-	-	-	✓	-	-	✓	-	6 MC	1 ML	5 R	✓	-	✓	✓	23	S	TB/TL
View of Delft [10]	✓	-	-	-	✓	-	-	-	-	1 SC	1 ML	1 R	✓	-	✓	-	13	-	TB
Waymo Open [16]	✓	-	-	-	✓	-	✓	-	-	5 MC	5 ML	-	-	✓	✓	✓	4	S	-
A2D2 [13]	✓	✓	✓	-	✓	-	✓	-	-	6 MC	5 ML	-	-	-	✓	✓	38	-	TB/TL
Cityscape (and 3D) [14]	✓	-	-	-	✓	-	✓	-	-	1 MC	-	-	-	-	✓	✓	30	-	-
Radiate [4]	✓	-	✓	✓	✓	✓	✓	✓	✓	1 SC	1 ML	1 R (360°)	✓	✓	(✓)	-	8	-	TB
Pixset [8]	✓	✓	-	-	✓	-	(10%)	-	-	3 MC	1 ML/1 SSL	1 R	✓	-	✓	-	20	PTP	TB
aiMotive [15]	✓	-	✓	-	✓	-	(4%)	-	-	5 MC	1 ML	2 R	✓ (GNSS)	✓	✓	-	14	-	-
Rain Wcity [5]	✓	-	-	-	✓	-	1 MC	-	-	1 MC	-	-	-	-	-	-	-	-	-
EU Long-term [6]	✓	✓	-	-	✓	-	-	-	✓	2 SC/2 MC	3 ML/1 SSL	1 R	✓ (GNSS)	-	-	-	-	H/S	TB
Zenseact [3]	✓	✓	✓	-	✓ (80%)	✓ (19%)	✓ (16%)	✓ (2%)	✓ (2%)	1 MC	3 ML	-	✓ (GNSS)	✓	✓	✓	15	-	-

3 Fulfilling the Requirements of Sensor Fusion

3.1 Sensors and Sensors Placement



(a) Prototyping vehicle



(b) Sensor Placement Architecture

Figure 1: Perception System

Expleo Germany GmbH has been developing a prototyping vehicle and corresponding functionalities for highly automated and autonomous driving functions like AVP and

Autonomous Valet Driving System (AVDS) in the last few years. The newly enhanced sensor setup integrated into Expleo’s vehicle collects data in real-world scenarios. It consists of two sub-systems: the AD setup and the reference setup. The autonomous driving (AD) setup includes a set of sensors that can potentially be integrated into a production vehicle, and it’s composed of three solid-state LiDARs (SSL), Robosense M1, two stereo cameras (SC), four short-range radars (SRR), and two long-range radars (LRR). Those sensors have been integrated into the vehicle chassis as shown in Figure 1b. The reference sensor setup consists of three mechanical LiDARs (one Velodyne VLP-32 and two Robosense BlackPearls(BP)) positioned atop the vehicle to model the prototyping vehicle’s entire surroundings, as depicted in Figure 1a.

3.2 Sensor Calibration

For our prototype vehicle sensor setup, manual calibration can prove challenging due to some sensors being embedded in the vehicle chassis, making their centers invisible. To overcome this, we have explored the possibility of estimating the transformation between sensors using the data they generate. In this study, our focus is on target-based multi-sensor calibration. It is a challenging task to extract environmental features with radars that are comparable to LiDARs or cameras as well as obtaining an accurate trajectory with radars. In [17], a multi-sensor calibrator (MSCT) was proposed to calibrate sensors similar to our setup, including LiDAR, camera, and radar. The Target comprises of four rings in the center of a rectangular board alongside a radar reflector positioned in the middle behind the board. The 3D positions of the four rings’ centers are estimated as a pattern from the LiDAR and camera data in addition to the corner reflector from the radar data. The optimization method is further refined by incorporating a loop closure constraint between the sensors.

However, we discovered while testing this calibration procedure on our sensor setup, the patterns were not detectable by the BP LiDAR. This is because the implemented detection procedure is based on the horizontal line scans that reflect the rings on the board. Consequently, we chose an alternative solution that is not dependent on the LiDAR’s point cloud structure. We replaced the rings board (RB) with a chessboard (CB), which is a widely used target for camera-2-LiDAR calibration [18]. For the detection of chessboard edges using cameras, we implemented the commonly-used pattern detection method [19]. To detect chessboard edges with LiDAR, we first estimated the point clouds that belong to the chessboard using the RANSAC algorithm. We then used the convex hull algorithm to extract the border of the chessboard, and based on that, used a minimum rectangle fitting algorithm to extract the edges of the chessboard.

For comparing and validating the various calibration solutions tested, we constructed a test bench as illustrated in Figure 2. To accurately estimate the sensor center, we designed customized mountings to digitally measure the distance between the sensor center and a reference point on the mountings. We manually measured the distances between these designated reference points for each sensor to calculate the transformation matrix between the sensors.



Figure 2: Test bench for testing different sensor calibration tools

Table 2: Results of multiple calibration tools on different sensor-to-sensor constellations

Sensor-to-Sensor	Calibration tool	Rotation Error (deg)	Translation error (cm)
VLP-32 to camera	MSCT Tool with CB	[0.83 0.52 0.74]	[2.08 1.31 3.94]
	MSCT Tool with RB	[0.3 0.72 1.33]	[2.09 1.38 3.9]
	Matlab tool	[0.15 0.21 2.1]	[2.02 3.17 0.19]
VLP-32 to BP	MSCT Tool	[1.51 0.25 1.164]	[2.62 0.67 1.39]
M1 to LRR	MSCT Tool	[0.074 1.49 1.41]	[1.3 3.08 5.39]

Table 2 displays the outcomes of our calibration method for various sensor-to-sensor configurations examined on the constructed test bench. We also evaluated the VLP-32-to-camera calibration with alternate methodologies like the MSCT Tool with RB, and the Lidar and Camera Calibration tool from Matlab. We observed similar results for this sensor constellation. However, in general, the calibration results differ with an average 5.5 cm translational and 2-degree rotational error from manual calibration. This can have a negative impact on the accuracy of sensor fusion applications. It is also worth mentioning that the positioning of the sensors on the test bench differs from the prototyping vehicle setup and can lead to diverse calibration results.

3.3 Time Synchronisation

As mentioned in section 2, time synchronization plays a crucial role in sensor fusion, particularly for autonomous driving applications. A time offset between different sensors can lead to significant positional errors when observing the same object, especially in highly dynamic conditions. As outlined in [20], the delay in message delivery can be divided into three components: the delay from the sender (sensors), the delay from propagation, and the delay from the receiver (data logger).

Conversely, the time delay from the receiver is non-deterministic due to potential operating system overheads. When the data logger approaches its data processing capacity limit, a conventional solution like that described in [21], which uses the data logger to assign timestamps to different sensors, struggles to assign timestamps reliably. This challenge prompted the development of our time synchronization framework, specifically designed to address time delays in message delivery (see figure 3).

The propagation delay is deterministic and relies on the distance between the sender and the receiver. In our use case, the delay from the sender (the sensors) is also deterministic because it is implemented on a hardware level. On the other hand, the time delay from the receiver is non-deterministic due to possible overheads from the operating

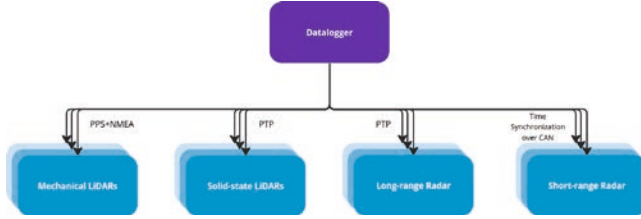


Figure 3: Overview of the created time synchronization framework

system. When the data logger nears its data processing limit, a conventional solution like the one described in [21], which uses the data logger to assign timestamps to different sensors, struggles to assign timestamps reliably. To address this issue, we developed a time synchronization framework (refer to figure 3) that is specifically tailored to handle message delivery delays.

In our framework, mechanical LiDARs (e.g. VLP-32C and BP) use the GPS time synchronization method [22, 23]. However, in underground parking garages, there is no GPS signal. As a solution, we propose an approach that simulates GPS signals via our data logger’s clock. We simulate the GPS signal for mechanical LiDARs by producing the Pulse Per Second (PPS) and National Marine Electronics Association (NMEA) sentence, following a previous study [24].

Precision Time Protocol (PTP) is used for time synchronization of solid-state LiDARs (M1) and LRRs. PTP, introduced in 2002 as a method for synchronizing clocks in distributed systems (Eidson, 2002), is used in conjunction with the data logger as the master clock and the solid-state LiDARs and long-range radar as the slave clock.

The SRRs exclusively support time synchronization over the Controller Area Network (CAN). We adhere to the specifications outlined in the Automotive Open System Architecture (AUTOSAR) [25] to implement this time synchronization protocol.

Table 3: Time Synchronization Framework comparison

Sensors	Typical Solution STD(ms)	Our STD(ms)
VLP-32C	0.15	0.02
BlackPearl	0.4	0.05
M1	0.1	0.03
SRR	4.4	1.2
LRR	3.6	0.6

We compare our framework with the typical solution, which assigns the arrival time to the message when the message arrives at the data logger [21]. We use the standard deviation of the sensor clock with the time of arrivals, which is an evaluation metric used in [24], to evaluate our time synchronization protocol.

As shown in table 3, we collected 3000 messages and measured the stability of the packet-to-packet period by calculating the standard deviation. Our framework shows more stability because the operating system overheads will not affect the time synchronization performance.

4 Methodology

4.1 Longterm AVP Data Generation in Controlled Environment

Currently, available open datasets for AD's perception module are reliant on data collected in the real world where critical situations and adverse weather conditions are infrequent. Additionally, research is confined to using sensor setups that may not satisfy the demands of new sensor arrangements, and where the collected data needs precise annotations. In most modern datasets, LiDAR sensors are frequently utilized as a reference owing to their high accuracy. However, during adverse weather conditions like foggy weather, LiDAR signal performance can be significantly compromised, which negatively impacts the quality of the annotation process. To resolve these concerns, we present a methodology for collecting and annotating data for vehicle ego localization as well as object detection and prediction in an LAMP context, specifically in regard to pedestrian-related scenarios, which is the most common Vulnerable Road User present in LAMP scenarios. Considering the ISAFE indoor testing facility, it is feasible to generate various scenarios involving harsh weather conditions, such as varying rain and fog intensities, as well as controlled lighting conditions. In the available test area, one can simulate different real-world scenarios. Parking situations are also taken into account for data collection. A brief outline of the potential weather conditions that could be incorporated into the data is presented in figure 4.



Figure 4: Scenario variety for pedestrian detection at ISAFE.

Each condition presents distinct challenges for various sensors. For example, cameras face visibility issues while the radar point cloud may have ghost objects when multiple objects are in close proximity. With different combinations collected with different objects, it is possible to either improve or develop new sensor fusion algorithms to achieve a robust performance under different conditions. Additionally, it can serve as a complementary source to the existing state-of-the-art datasets.

4.2 Reference Systems

Reference systems are essential in order to evaluate the accuracy of the developed perception functions such as Object Detection and Classification or Vehicle ego-localization.

4.2.1 LiDAR-based Reference Localization System

For testing in GPS-deprived environments, utilizing localization systems that are integrated into the infrastructure yields an accurate solution, such as the IPS system (discussed in section 4.2.2) or MOCAP systems. Nevertheless, conducting real-world testing in various locations incurs significant overhead for integrating and calibrating these systems into the infrastructure. In our use case, we chose to implement a LiDAR-based localization system, as it provides an accurate and robust solution.

We examined various benchmarks for LiDAR-based localization and selected two solutions: HDL-Graph-SLAM [26] and LEGO_LOAM [27]

HDL-Graph-SLAM generates and optimizes a graph where nodes indicate the sensor positions, and the edges between these nodes represent the odometry constraints (the relative pose between nodes), generated in this case by a point cloud scan matching algorithm, such as GICP [28]. The graph is optimized so that the error function between constraints and positions is minimized. With the help of **floor detection** and **loop-closure**, the graph is optimized and a more precise map of the environment is produced. the LiDAR frames are afterwards reprocessed through a scan-matching algorithm, where these LiDAR frames are registered with the generated map to estimate the final trajectory.

LEGO-LOAM is a LiDAR-based SLAM algorithm based on feature matching. This algorithm discretizes the search space according to the LiDAR point cloud structure, to extract 3D points that belong to edges and surfaces. It estimates the spatial transformation between two consecutive frames by finding correspondences between extracted edges and surface features. In addition, a LiDAR mapping module is used to refine the pose transformation. It matches the features extracted from the latest frame with all the features extracted from the oldest frames stored in a map, which is further optimized by a loop closure detection method.

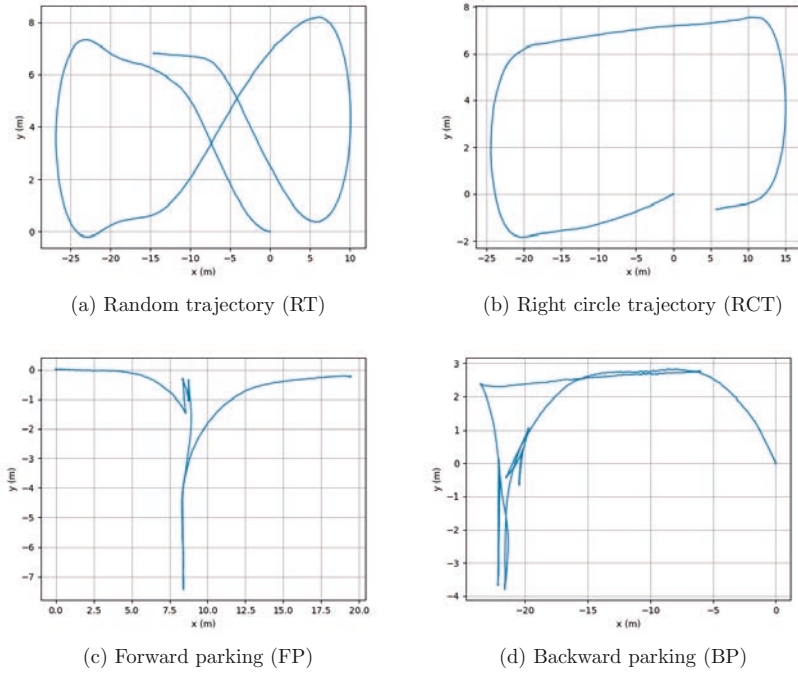


Figure 5: Different tested parking scenarios trajectories

Table 4: LiDAR-based localization results on different trajectories

Method	Evaluation Metrics (m)	RT	RCT	FP	BP
LEGO-LOAM	RMSE	0.07	0.162	0.077	0.074
	MAX Error	0.1433	0.249	0.149	0.115
HDL-GS	RMSE	0.058	0.144	0.086	0.094
	MAX Error	0.169	0.252	0.159	0.191

Figure 5, represents an example of relevant trajectories that we used for testing our Reference localization system, and the table 4 represents the results for these different trajectories. We observe that while the root mean square error (RMSE) results for most of the trajectories are below 10 cm, the maximal error reaches 0.24 cm. This presents a challenge since a reference system should always maintain accuracy across all scenarios. Thus, further improvement of this system is necessary.

4.2.2 Data Annotation Methods

Annotating data is expensive and time-consuming, making evaluating and testing perception systems a challenge. We propose an active learning-based labeling tool and an indoor positioning system-based annotation approach to address this challenge.

Active Learning Based Semi-Auto Labeling Tool for Semantic Segmentation

At Expleo, we created a semi-automatic labeling tool for semantic segmentation that employs active learning to reduce the manual effort of human annotation. Notably, we have incorporated the One-vs-All (OVA) method into active learning and found that it enhances diversity for active selection, leading to improved segmentation accuracy [29]. Using uncertainty, the Semi-Auto labeling tool suggests potential candidate points of an image and asks the human annotators for the true label. The human annotation is utilized to retrain the neural network. According to [29], this semi-automatic labeling tool achieves 96% performance in fully supervised learning with 100 pixels per image (0.06% of the entire dataset) on CityScapes [14]. We incorporate this semi-automatic labeling tool in our reference system pipeline to facilitate the prompt and efficient generation of image segmentation labels by human annotators.

IPS Reference Based Annotation For the annotation process, we use the IPS that allows for accurate labeling regardless of weather conditions. Equipped with a receiver antenna, the positioning system can provide precise information, including coordinates, velocity, and heading angle of the object of interest. To effectively develop and test detection methods, reliable ground truth data is crucial, a necessity fulfilled by this positioning system. This system offers significant advantages in the domain of pedestrian and vehicle detection tasks that are based on vision-based sensors like camera, LiDAR, and radar sensors. A notable benefit is its ability to maintain accuracy even in inclement weather conditions, such as rain and fog, where other sensors may produce unreliable or incomplete data. Furthermore, it can generate both 3D and 2D bounding boxes for training object detection models using images and point clouds. See Figure 6 for an illustration of the data collection setup and annotation process.

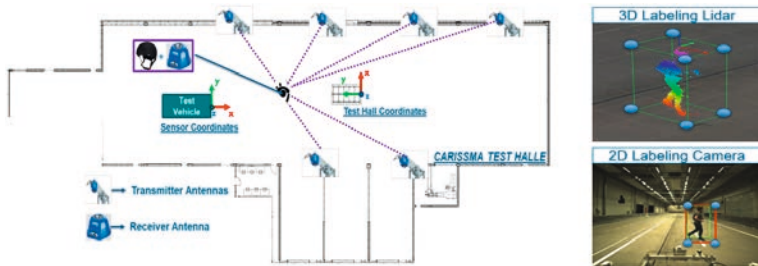


Figure 6: Overview of the annotation framework at ISAFE.

To conduct the annotation process, first, create a 3D bounding box that encloses the entire body of the pedestrian across all data frames in a specific experiment. This can

be accomplished using IPS measurements for each coordinate (X, Y, and Z in meters). The reference point of the Z-axis remains at ground level. Using this data, the height (corresponding to the Z dimension of the 3D bounding box) is calculated. Subsequently, predetermined values for the X and Y dimensions are applied. In a second step, the corners of the 3D bounding box project onto the target image (2D space) employing the camera parameters and calibration values obtained from Chapter 3.3. The minimum and maximum values within the image boundaries establish the final 2D bounding box, as illustrated in figure 6. This method can be used regardless of weather conditions. However, as depicted in figure 6, the resulting bounding box (indicated in red), may be larger than the actual size of the object depending on the viewpoint and perspective. Furthermore, this technique can produce complete body bounding boxes even when the pedestrian is partially or entirely obstructed.

Complementarity of Annotation Methods Considering our context, the presented methods can complement each other to offer a more time-efficient approach to data generation. In real-world scenarios where no positioning system is available, the active learning-based method could prove useful for annotation generation. In contrast, in edge cases with harsh weather, the labeling method aided by the IPS presents a viable option. Hence, both methods offer a viable approach to gather supplementary scenarios, both real and in a controlled setting, to create a comprehensive dataset encompassing all possible test cases with variations in weather and scenarios to facilitate complete training and assessment of algorithms.

5 Conclusion

The evaluation of various pre-existing datasets for the development of a long-range automated vehicle revealed open issues in their application due to several factors, including localization in an indoor environment and the acquisition and annotation of adverse weather conditions. We propose a method of generating complementary data, taking advantage of the closed indoor environment found at ISAFE, and emphasize the benefits of an IPS reference-based annotation process. Table 5 shows our method and its possible features in regard to state-of-the-art datasets mentioned in chapter 2. The presented sensor setup is composed of a relatively diverse set of sensors in comparison with state-of-the-art datasets. For multi-sensor calibration, we implemented a solution adequate for our sensor setup. Through a test bench, we showed that the reached accuracy is still not enough in comparison with manual calibration and this can impact using sensor fusion for different use cases. For time synchronization, we showed that our framework provides more stability. Furthermore, we presented our reference systems. In future works, we will further investigate multi-LiDAR-based localization to enhance the accuracy of our reference localization system. This will give us more insight into the needed accuracy for sensor calibration. For the latter, we plan to investigate target-less sensor calibration methods.

Table 5: State-of-the-art datasets. Abbreviations scenarios: (U)urban, (SU)suburban, (H)highway and (P)parking; conditions: (G)good, (N)night, (R)rain, (F)fog and (S)snow; sensors: (M)mono/(S)stereo (C)camera, (M)mechanical/(SS)solid-state (L)LiDAR and (R)radar; synchronization (Sync.): (H)hardware and (S)software; calibration: (M>manual, (TL)target-based and (TB)target-less; (-)no information available or not available in the dataset.

Dataset	Scenarios				Conditions				Sensors			GPS/IMU	Annotation			Classes	Sync.	Calib.
	U	SU	H	P	G	N	R	F	S	C	L	R	2D	3D	Sem.			
KITTI [9]	✓	✓	✓	-	✓	-	-	-	-	✓	✓	-	✓	-	✓	28	H/S	TB
nuScene [2]	✓	-	-	-	✓	✓	✓	✓	-	✓	✓	✓	✓	✓	✓	23	H/S	TB/TL
.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.
.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.
Our Method	scenario selectable in regard to possibilities of a testing area				✓	✓	✓	✓	-	2 SC	3 ML/ 3 SSL	6 R	IPS	IPS for moving objects 2D & 3D			S	TB

6 Acknowledgements

The research leading to these results is funded by the Bavarian Ministry of Economic Affairs, Energy and Technology and by Expleo Germany GmbH within the project: "SAFE2P - Increase Safety for Pedestrians and Parking" (DIK-350).

References

- [1] Muhammad Khalid, Kezhi Wang, Nauman Aslam, Yue Cao, Naveed Ahmad, and Muhammad Khurram Khan. From smart parking towards autonomous valet parking: A survey, challenges and future works. *Journal of Network and Computer Applications*, 175:102935, 2021.
- [2] Holger Caesar, Varun Bankiti, Alex H. Lang, Sourabh Vora, Venice Erin Liong, Qiang Xu, Anush Krishnan, Yu Pan, Giancarlo Baldan, and Oscar Beijbom. nusenes: A multimodal dataset for autonomous driving. *arXiv preprint arXiv:1903.11027*, 2019.
- [3] Mina Alibeigi, William Ljungbergh, Adam Tonderski, Georg Hess, Adam Lilja, Carl Lindstrom, Daria Motorniuk, Junsheng Fu, Jenny Widahl, and Christoffer Petersson. Zenseact open dataset: A large-scale and diverse multimodal dataset for autonomous driving, 2023.
- [4] Marcel Sheeny, Emanuele De Pellegrin, Saptarshi Mukherjee, Alireza Ahrabian, Sen Wang, and Andrew Wallace. Radiate: A radar dataset for automotive perception. *arXiv preprint arXiv:2010.09076*, 2021.
- [5] Xian Zhong, Shidong Tu, Xianzheng Ma, Kui Jiang, Wenxin Huang, and Zheng Wang. Rainy wecity: A real rainfall dataset with diverse conditions for semantic driving scene understanding. In Lud De Raedt, editor, *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, pages 1743–1749. International Joint Conferences on Artificial Intelligence Organization, 7 2022. Main Track.

- [6] Zhi Yan, Li Sun, Tomas Krajník, and Yassine Ruichek. EU long-term dataset with multiple sensors for autonomous driving. In *Proceedings of the 2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2020.
- [7] De Jong Yeong, Gustavo Velasco-Hernandez, John Barry, and Joseph Walsh. Sensor and sensor fusion technology in autonomous vehicles: A review. *Sensors*, 21(6):2140, 2021.
- [8] Jean-Luc D’Áziel, Pierre Merriaux, Francis Tremblay, Dave Lessard, Dominique Plourde, Julien Stanguennec, Pierre Goulet, and Pierre Olivier. Pixset : An opportunity for 3d computer vision to go beyond point clouds with a full-waveform lidar dataset, 2021.
- [9] Andreas Geiger, Philip Lenz, Christoph Stiller, and Raquel Urtasun. Vision meets robotics: The KITTI dataset. *International Journal of Robotics Research*, 32(11):1231 – 1237, September 2013.
- [10] Andras Palffy, Ewoud Pool, Srimannarayana Baratam, Julian F. P. Kooij, and Darius M. Gavrilă. Multi-class road user detection with 3+1d radar in the view-of-delft dataset. *IEEE Robotics and Automation Letters*, 7(2):4961–4968, 2022.
- [11] Tuomas Välimäki, Bharath Garigipati, and Reza Ghabcheloo. Motion-based extrinsic sensor-to-sensor calibration: Effect of reference frame selection for new and existing methods. *Sensors*, 23(7):3740, 2023.
- [12] Chongjian Yuan, Xiyuan Liu, Xiaoping Hong, and Fu Zhang. Pixel-level extrinsic self calibration of high resolution lidar and camera in targetless environments. *IEEE Robotics and Automation Letters*, 6(4):7517–7524, 2021.
- [13] Jakob Geyer, Yohannes Kassahun, Mentar Mahmudi, Xavier Ricou, Rupesh Durgesh, Andrew S. Chung, Lorenz Hauswald, Viet Hoang Pham, Maximilian Mühlegg, Sebastian Dorn, Tiffany Fernandez, Martin Jänicke, Sudesh Mirashi, Chiragkumar Savani, Martin Sturm, Oleksandr Vorobiov, Martin Oelker, Sebastian Garreis, and Peter Schuberth. A2D2: Audi Autonomous Driving Dataset. 2020.
- [14] Marius Cordts, Mohamed Omran, Sebastian Ramos, Timo Rehfeld, Markus Enzweiler, Rodrigo Benenson, Uwe Franke, Stefan Roth, and Bernt Schiele. The cityscapes dataset for semantic urban scene understanding. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 3213–3223, 2016.
- [15] Tamas Matuszka, Ivan Barton, Ádám Butykai, Péter Hajas, Dávid Kiss, Domonkos Kovács, Sándor Kunsági-Máté, Péter Lengyel, Gábor Németh, Levente Pető, Dezső Ribli, Dávid Szeghy, Szabolcs Vajna, and Balint Viktor Varga. aimotive dataset: A multimodal dataset for robust autonomous driving with long-range perception. In *International Conference on Learning Representations 2023 Workshop on Scene Representations for Autonomous Driving*, 2023.
- [16] Pei Sun, Henrik Kretzschmar, Xerxes Dotiwalla, Aurelien Chouard, Vijaysai Patnaik, Paul Tsui, James Guo, Yin Zhou, Yuning Chai, Benjamin Caine, Vijay Vasudevan,

- Wei Han, Jiquan Ngiam, Hang Zhao, Aleksei Timofeev, Scott Ettinger, Maxim Kriukon, Amy Gao, Aditya Joshi, Yu Zhang, Jonathon Shlens, Zhifeng Chen, and Dragomir Anguelov. Scalability in perception for autonomous driving: Waymo open dataset. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2020.
- [17] Joris Domhof, Julian FP Kooij, and Darius M Gavrilă. A joint extrinsic calibration tool for radar, camera and lidar. *IEEE Transactions on Intelligent Vehicles*, 6(3):571–582, 2021.
- [18] Lipu Zhou, Zimo Li, and Michael Kaess. Automatic extrinsic calibration of a camera and a 3d lidar using line and plane correspondences. In *2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 5562–5569. IEEE, 2018.
- [19] Andreas Geiger, Frank Moosmann, Ömer Car, and Bernhard Schuster. Automatic camera and range sensor calibration using a single shot. In *2012 IEEE international conference on robotics and automation*, pages 3936–3943. IEEE, 2012.
- [20] Ying Weng and Yiming Zhang. A survey of secure time synchronization. *Applied Sciences*, 13(6):3923, 2023.
- [21] Puck lidar sensor (vlp-16). velodyne lidar. <https://velodynelidar.com/products/puck/dc>. Accessed: 2023-08-17.
- [22] Robosense blackpearl: Hemispherical super wide fov, short-range blind spot lidar. <https://www.robosense.ai/en/rslidar/RS-Bpearl>. Accessed: 2023-08-17.
- [23] Ultra puck: The high-density, long-range image generated by the ultra puck makes it an industry favorite for robotics, mapping, security, driver assistance, and autonomous navigation. <https://velodynelidar.com/products/ultra-puck/downloads>. Accessed: 2023-08-17.
- [24] Marsel Faizullin, Anastasiia Kornilova, and Gonzalo Ferrer. Open-source lidar time synchronization system by mimicking gnss-clock. In *2022 IEEE International Symposium on Precision Clock Synchronization for Measurement, Control, and Communication (ISPCS)*, pages 1–5. IEEE, 2022.
- [25] Specification of time synchronization over can autosar cp r21-11. https://www.autosar.org/fileadmin/standards/R21-11/CP/AUTOSAR_SWS_TimeSyncOverCAN.pdf. Accessed : 2023 – 08 – 17.
- [26] Kenji Koide, Jun Miura, and Emanuele Menegatti. A portable three-dimensional lidar-based system for long-term and wide-area people behavior measurement. *International Journal of Advanced Robotic Systems*, 16, 02 2019.
- [27] Ji Zhang and Sanjiv Singh. Low-drift and real-time lidar odometry and mapping. *Autonomous Robots*, 41(2):401 – 416, February 2017.
- [28] Aleksandr Segal, Dirk Haehnel, and Sebastian Thrun. Generalized-icp. In *Robotics: science and systems*, volume 2, page 435. Seattle, WA, 2009.

- [29] Jing Gu. One-vs-all semi-auto labeling tool for semantic segmentation in autonomous driving. July 2023.

Ein Ansatz zur automatisierten Erstellung von Trainingsdaten unter Verwendung von HD-Karten und Mehrfachbefahrungen

Frank Bieder* Haohao Hu† Johannes Schantz‡ Oguzahn Kirik§
 Florian Ries¶, Martin Haueis|| und Christoph Stiller**

Zusammenfassung:

In diesem Beitrag wird ein Gesamtsystem zur skalierbaren Erstellung von Trainingsdaten für das maschinelle Lernen im Kontext des automatisierten Fahrens vorgestellt. Unter Verwendung einer hochgenauen Lokalisierung in verifizierten HD-Karten und mithilfe von Sensor-Projektionsmodellen können semantische Kartenmerkmale in die verschiedenen Sensordomänen rückprojiziert werden. Hierdurch entsteht eine Korrespondenz von statischen Kartenmerkmalen und den entsprechenden Messaufnahmen, welche für das Training von neuronalen Netzen verwendet werden können. Durch Mehrfachbefahrungen und eine Multi-Drive-Kartierung kann die Anzahl der Trainingsdaten nach Belieben skaliert werden. Durch die Methode lassen sich statische Kartenmerkmale wie Fahrbahnlinsen und Straßenmarkierungen automatisiert annotieren. Es werden Verfahren vorgestellt, um sowohl mit statischen Verdeckungen, wie Pfosten, als auch mit dynamischen Verdeckungen, wie Fahrzeuge oder Fußgänger, umzugehen.

Für die Umsetzung der Methode wurde mit dem Versuchsfahrzeug BerthaOne ein Datensatz in den Städten Karlsruhe und Sindelfingen aufgenommen. Dieser besteht aus mindestens 20 Kilometern Strecke und umfasst 4 Befahrungen pro Strecke. Anhand dieses Datensatzes wird das System und die Qualität der erzeugten Trainingsdaten qualitativ evaluiert.

Schlüsselwörter: Automatisiertes Fahren, skalierbare Erstellung von Trainingsdaten, Deep Learning, HD Karten, Mehrfachbefahrungen

1 Einleitung

Für den sicheren und erfolgreichen Einsatz von hoch-automatisierten Fahrfunktionen wird ein umfassendes Verständnis der Umgebung benötigt. Hierbei werden immer mehr Aufgaben des Umgebungsverständnisses durch tiefe neuronale Netze gelöst. Dadurch ist es möglich, die Umgebung in Echtzeit einzuschätzen und entsprechende Handlungen daraus abzuleiten. Für die performante Anwendung von tiefen neuronalen Netzen werden große

*Frank Bieder ist Doktorand am FZI Forschungszentrum Informatik, E-mail: bieder@fzi.de.

†Haohao Hu ist Doktorand am MRT Institut für Mess- und Regelungstechnik, KIT.

‡Johannes Schantz ist CTO von enabl Technologies GmbH, zuvor Masterand am MRT, KIT.

§Oguzahn Kirik ist Doktorand an der TU Berlin, zuvor Masterand am MRT, KIT.

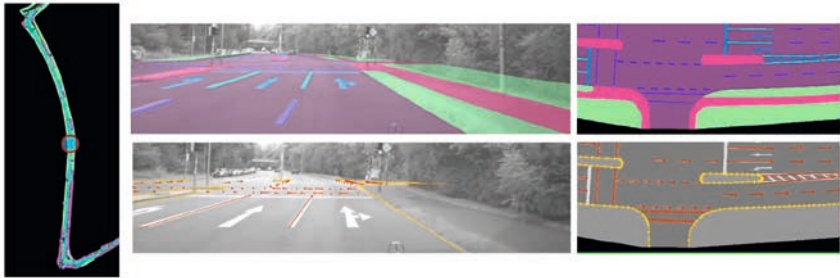
¶Florian Ries ist Entwicklungsingenieur bei der Mercedes-Benz Group AG.

||Martin Haueis ist Teamleiter bei der Mercedes-Benz Group AG.

**Christoph Stiller ist Professor am MRT Institut für Mess- und Regelungstechnik, KIT.



(a) Beispiel für eine der verwendeten HD-Karten in der Vogelperspektive. Der Kartenabschnitt zeigt eine Kreuzung in Karlsruhe und ist in drei verschiedenen Zoom-Stufen dargestellt. Für die Luft- und Satellitenbilder wurde Esri World Imagery [1] verwendet.



(b) Generierung von Trainingsdaten. Gegeben einer beliebigen 6D-Pose aus einer Befahrung (links) können Kartenelemente in die verschiedenen Sensordomänen rückprojiziert werden. In dem Beispiel werden Trainingsdaten für semantische Segmentierung und für die Detektion von instanzbasierten Liniensegmenten generiert - jeweils für die Frontkamera- und Vogelperspektive.

Abbildung 1: Visualisierung der Datenbasis und exemplarische Darstellung der vorgestellten Methode für verschiedene Merkmalstypen und Sensordomänen.

und diverse Datensätze benötigt, welche möglichst genau auf das vorhandene Sensor-Setup und die Zieldomäne passen. Das Aufnehmen und Erstellen dieser Datensätze ist ein aufwendiger und kosten-intensiver Prozess. Hinzu kommt, dass sich durch den Wechsel des Sensor-Setups oder der Zieldomäne eine Lücke zwischen der Domäne der Sensordaten und der Domäne des Datensatzes bildet. Wird diese Lücke zu groß, muss der Prozess der Datensatzerstellung wiederholt werden. Die vorliegende Arbeit stellt eine Methode vor, um Datensätze für das maschinelle Lernen effizient und skalierbar durch die Verwendung von HD-Karten und einer Multi-Drive-Kartierungs-Methode automatisiert zu annotieren. Dabei werden Karten-Merkmale der HD-Karte in die Sensoraufnahmen rückprojiziert und dadurch die Ground Truth für die jeweiligen Sensoraufnahmen erzeugt. Die statischen Informationen einer Szene können aus verschiedenen Perspektiven erfasst und zur mehrfachen Annotation der Sensordaten verwendet werden. Durch die Datenrepräsentation in Form einer Karte reicht es also aus, Informationen einmal zu annotieren, um Trainingsdaten für unterschiedliche Sensoren und unterschiedliche Verkehrssituationen zu erstellen. Abbildung 1 stellt die Methode für verschiedene Merkmalstypen und Sensordomänen dar.



Abbildung 2: Auszüge der kartierten Datenbasis. Links und in der Mitte sind zwei Strecken aus Karlsruhe. Die rechte Strecke wurde in Sindelfingen kartiert. Im Datensatz sind unterschiedliche Szenen von verschiedener Komplexität wie zum Beispiel ein- oder mehrspurige Straßen im städtischen Umfeld mit Fahrradspuren und Fußgängerüberwegen und Kreuzungen enthalten, aber auch am Stadtrand gelegene Straßen, welche Landstraßen ähneln.

2 Stand der Technik

Die manuelle Annotation eines Datensatzes ist sehr zeitintensiv. Dieser hohe zeitliche Aufwand steigt mit der zunehmenden Komplexität der zu annotierten Merkmale und der zunehmenden Größe der benötigten Daten für moderne Modelle des maschinellen Lernens. Im Folgenden wird ein Überblick über die verschiedenen Ansätze zur Erstellung von Datensätzen im Kontext des automatisierten Fahrens gegeben.

Für die pixel-weise semantische Annotation eines Bildes des Cityscapes Datensatzes wurden laut [2] im Durchschnitt 1,5 Stunden benötigt. Hierbei muss jedoch beachtet werden, dass dies die Qualitätskontrolle der Annotation beinhaltet und der Datensatz zu seiner Zeit neue Maßstäbe an Annotations-Qualität gesetzt hat. [3] verweist auch auf eine Annotationszeit von 60 Minuten pro Bild im manuellen Fall und schlägt daher ein Verfahren vor, welches 2D-Labels im 3D-Raum erzeugt. Hierbei wird der Labelvorgang mithilfe von Formprimitiven direkt in der 3D-Punktwolke durchgeführt. Die Informationen werden anschließend in die 2D-Domäne projiziert. In [8] wird ein Ansatz vorgestellt, welcher eine akkumulierten LiDAR Punktwolke semantisch annotiert. Die semantischen Informationen können anschließend sowohl auf einzelne LiDAR Punktwolken als auch auf Kamerabilder projiziert werden. Inzwischen gibt es viele weitere Ansätze, welche versuchen, die Annotation von Trainingsdaten teil-zuautomatisieren. Eine Möglichkeit hierfür ist die Verwendung von Methoden des maschinellen Lernens, um den Annotationsprozess zu unterstützen. [4] beschreibt ein Verfahren, bei welchem ein Mensch durch das ziehen einer Bounding Box oder das Verschieben von Polygonpunkten Bedingungen setzt, welche von einem Graph Neural Network interpretiert werden. Das Modell schlägt anschließend in jedem Iterationsschritt eine feinkörnigere Annotation vor. Dieser Vorgang wird so lange wiederholt, bis die erstellte Annotation von der annotierenden Person als korrekt akzeptiert wird. Synthetische Simulationen können auch dazu verwendet werden, mit verhältnismäßig wenig Aufwand eine große Menge an Trainingsdaten zu generieren. GTA5 [5] oder SYNTHIA [6] sind Beispiele für synthetische Datensätze. CARLA [7] hingegen bietet eine vollständige Simulationsumgebung um Trainingsdaten nach den eigenen Wünschen zu

generieren. Bei synthetischen Daten bleibt jedoch das Problem der Domän-Unterschiede zwischen realer Welt und synthetischer Welt bestehen, wodurch die Übertragbarkeit der Modelle in die reale Welt eingeschränkt ist. Neben [8] stellen auch die Datensätze Nuscenes [10], TuSimple und Argoverse [9] semantische Informationen in HD-Karten oder kartenähnlichen Datenstrukturen dar. Nuscenes liefert lediglich 2D Karten-Merkmale, wodurch eine Projektion in das Kamerabild ohne weiteres nicht möglich ist. TuSimple und Argoverse liefern 3D Kartenmerkmale sind allerdings auf Fahrstreifenränder konzentriert. Keiner dieser Datensätze liefert feingranularere Merkmale wie Straßenmarkierungen oder verfügt über mehrere Befahrungen mit einer Multi-Drive-Kartierung.

3 Automatisierte Erstellung von Trainingsdaten

Im Folgenden werden die Systemkomponenten vorgestellt, welche im Zusammenspiel die skalierbare Erstellung von Trainingsdaten ermöglichen.

Versuchsfahrzeug und Sensor-Setup

Die Messdaten wurden mit dem Versuchsfahrzeug BerthaOne [12] aufgezeichnet. Der Datensatz umfasst Daten, die von einem auf dem Dach montierten Velodyne Alpha Prime LiDAR, drei BlackFly PGE-50S5M-Kameras hinter der Front- und Heckscheibe und einem Ublox C94-M8P GNSS-Empfänger aufgenommen wurden. Der Alpha Prime LiDAR besitzt 128 Schichten und tätigt mit 10 Hz 360°-Aufnahmen. Die Kameras nehmen jeweils mit 10 Hz Bilder auf und der GNSS-Empfänger lieferte Messungen mit 5 Hz. Alle Sensoren wurden gemeinsam unter Verwendung der in [13] und [14] vorgestellten Ansätze intrinsisch und extrinsisch kalibriert. Der Datensatz besteht aus 7 Strecken mit einer Gesamtlänge von über 20 km, wobei 5 Strecken in Karlsruhe und 2 in Sindelfingen aufgenommen wurden. Drei Strecken sind in Abbildung 2 dargestellt. Hierbei wurden für alle Strecken mindestens 4 Befahrungen aufgenommen.

Multi-Drive-Kartierung und Mehrfachbefahrungen

In der vorliegenden Arbeit wird eine kamerabasierte SLAM-Methode [15] verwendet, welche in der Lage ist einen 6D-Posen-Graphen aus mehreren Befahrungen derselben Strecke zu erstellen. Jede weitere Befahrung erhöht die verfügbare Anzahl an potenzielle Trainingsdaten, einschließlich Kamerabildern und LiDAR-Punktwolken. Ein weiterer Vorteil der Multi-Drive-Kartierung ist die Verbesserung der Genauigkeit des resultierenden 6D-Posen-Graphs. Zusätzlich ermöglicht die Multi-Drive-Kartierung eine größere Bandbreite an Verkehrssituationen sowie Wetter- und Lichtverhältnissen derselben Szenerie. Dies erhöht die Diversität der Daten und kann damit die Robustheit der trainierten Modelle erhöhen. Je nach Anwendungsspezifikation und Charakteristik des urbanen Umfeldes kann die Anzahl der Befahrungen und damit die Skalierung der Trainingsdaten variiert werden. Hierbei ist zu beachten, dass der Mehrwert einer zusätzlichen Befahrung mit steigender Anzahl der Befahrungen abnimmt.

Manuelle Annotation von semantischen Kartenmerkmalen

Im Anschluss an die Multi-Drive-Kartierung kann unter Verwendung von Stereobildern ein Oberflächenabbild der Straße in der Vogelperspektive erstellt werden. Im Zuge dieser Arbeit wurde so anhand der ersten Befahrung einer Strecke eine Rasterkarte mit einer Auflösung von 2cm/Pixel erzeugt. Auf den dadurch entstandenen Rasterkarten sind Straßenmarkierungen und Straßenbegrenzungen sehr gut erkennbar und können manuell annotiert werden. Für weiteren Kontext können sowohl die entsprechenden Kamerabilder als auch Luft- oder Satellitenbilder hinzugezogen werden. Da die Höhe der Kartenmerkmale direkt aus den Stereobildern geschätzt wird, müssen 3D Merkmale auf der Straßenoberfläche lediglich in 2D annotiert werden. Im Zuge dieser Arbeit wurden 6 verschiedene Klassen von Kartenmerkmalen annotiert. Hierbei wurde zwischen drei Arten von Straßenmarkierungen und auch drei verschiedenen Flächen-Elementen unterschieden.

Kartenformat

Als Kartenformat wurde das Lanelet2-Framework [16] verwendet, welches speziell für den Einsatz im Bereich des automatisierten Fahrens entwickelt wurde. Es zeichnet sich unter Anderem durch die Einsatzmöglichkeiten in verschiedenen Anwendungen wie Lokalisierung oder Bewegungsplanung aus. Bei der Entwicklung von Lanelet2 wurden Genauigkeit, Vollständigkeit, Verifizierbarkeit und Erweiterbarkeit als Anforderungen an Karten für das automatisierte Fahren identifiziert und das Kartenformat entsprechend konzipiert. Die in dieser Arbeit erstellten HD-Karten enthalten semantische Informationen, welche im 3D-Raum als offene oder geschlossene Linienzüge definiert werden.

Repräsentation der Trainingsdaten

Aus dieser allgemeinen Kartenrepräsentation können Trainingsdaten in einer Vielzahl von Merkmals-Darstellungen erzeugt werden, einschließlich Linien und Flächen. Daher ist die Daten-Generierung nicht auf eine bestimmte Merkmalsrepräsentation limitiert. Es werden zwei Repräsentationen vorgestellt, welche in dieser Arbeit verwendet werden:

- Semantische Segmentierung: Die Zuordnung einer semantischen Klasse zu jedem Pixel bzw. jeder Zelle einer Rasterkarte.
- Detektionsverfahren für Liniensegmente nach [17]: Hierbei handelt es sich um ein instanzbasiertes Verfahren, welches Liniensegmente klassifiziert und lokalisiert.

Hierbei können die Kartenmerkmale in verschiedenen Sensordomänen rückprojiziert werden, um Trainingsdaten für verschiedene Sensortypen zu generieren. In dieser Arbeit werden zwei Sensor-Projektionsmodelle zur Erzeugung von Trainingsdaten verwendet, Frontview-Kamera und die Birds-Eye-View(BEV)-Domäne für Lidar-Rasterkarten.

Verdeckungen bei der Rückprojektion von Sensordaten

Bei der Erzeugung der Ground Truth für Kamerabilder besteht eine weitere wichtige Aufgabe darin Verdeckungen zu erkennen. Dies sollte verhindern, dass bei der Rückprojektion von Kartenmerkmalen in die Sensordaten nicht-sichtbare Kartenmerkmale als Elemente der Ground Truth gelabelt werden. Hierbei kann zwischen zwei Arten der Verdeckung unterschieden werden.

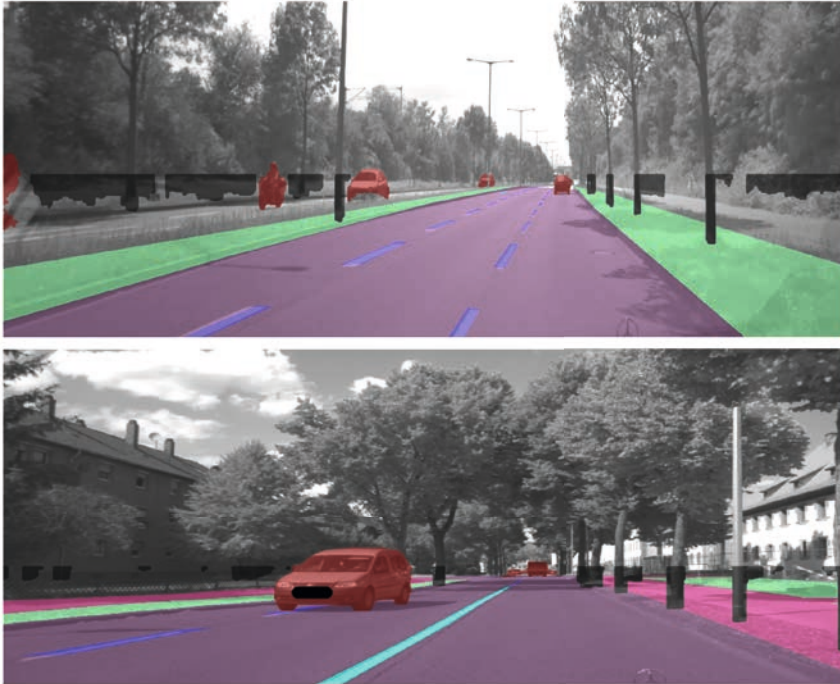


Abbildung 3: Qualitative Evaluation der automatisch generierten Trainingsdaten in Hinblick auf dynamische und statische Verdeckungen für den Fall der semantischen Segmentierung. In den Beispielen wurden alle statischen Klassen direkt aus einer HD-Karte rückprojiziert. Hingegen sind die dynamischen Objektklassen von einem neuronalen Netz geschätzt, um dynamische Verdeckungen zu identifizieren. Die grauen Flächen wurden als statische Verdeckungen mithilfe einer Stereo-Kamera identifiziert.

- **Dynamische Verdeckung:** Als dynamische Verdeckung werden sowohl Personen als auch Fortbewegungsmittel wie LKWs, Busse, Züge, Fahrräder, Motorräder und Autos bezeichnet.
- **Statische Verdeckung:** Als statische Verdeckungen werden unbewegliche Objekte bezeichnet, welche den Blick auf ein rückprojiziertes Kartenmerkmal verdecken können.

Da in beiden Fällen der Ansatz auf Sensordaten des aktuellen Zeitpunktes basiert, können nicht erkannte dynamische Verdeckungen auch durch den Ansatz für statische Verdeckung identifiziert werden. Allerdings ist der Ansatz für dynamische Verdeckungen besonders effektiv. Hierbei wird auf öffentliche Datensätze für semantische Segmentierung zurückgegriffen (Cityscapes [2], Mapillary [19]), welche Annotationen für die oben

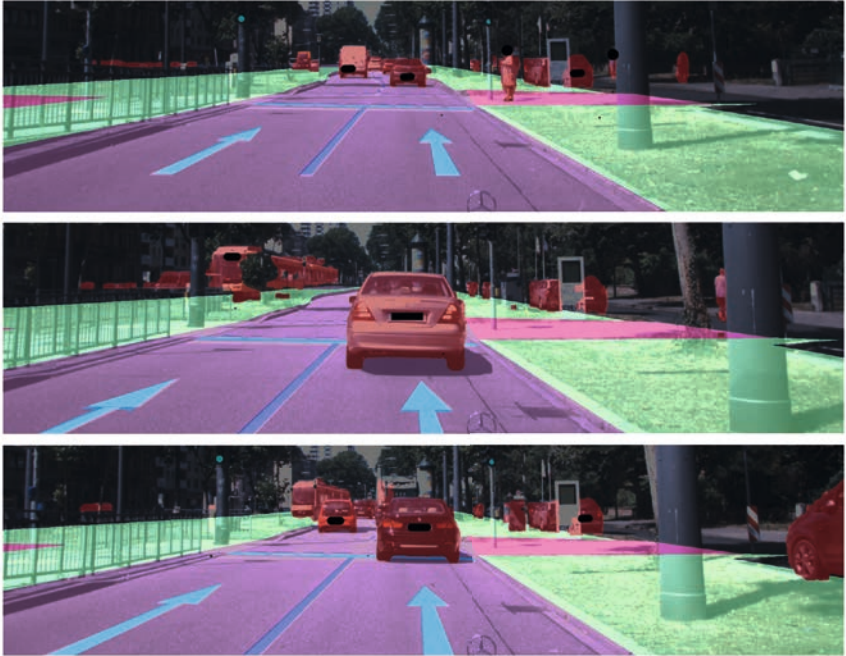


Abbildung 4: Qualitative Evaluation der Güte der Multi-Drive-Kartierung für den Fall der semantischen Segmentierung. Die Aufnahmen stammen aus drei verschiedenen Befahrungen derselben Strecke und wurden zu unterschiedlichen Zeitpunkten aufgenommen. Die Kartenmerkmale einer HD-Karte können in beliebig viele Befahrungen rückprojiziert werden, sofern es die Güte des 6D-Posen-Graphen zulässt.

genannten dynamischen Objektklassen besitzen. Unter Verwendung von großen, nicht-online fähigen Netzarchitekturen, wie beispielsweise SeamSeg: Seamless Scene Segmentation [18], werden diejenigen Pixel identifiziert, welche zu einem dynamischen Objekt gehören. Hieraus werden binäre Masken generiert, welche über die Annotation von rückprojizierten Kartenelementen entscheiden. Anschließend kann anhand von Tiefendaten aus der Stereokamera abgeschätzt werden, ob eine statische Verdeckung vorliegt. Hierbei wird nach homogenen Flächen im Tiefenbild gesucht, welche sich in der Sichtlinie eines Kartenobjektes befinden und eine deutlich geringere Distanz zu der Kamera aufweisen als das Kartenelement vermuten lässt. Hierbei können beispielsweise Pfosten, Bäume oder Häuserreihen an Kreuzungen erkannt werden und die Rückprojektion entsprechend angepasst werden. Dieses Verfahren ist stark von der Qualität des Tiefenbildes abhängig und hat sich als weniger robust als das Verfahren für dynamische Verdeckungen erwiesen.

4 Fazit

In diesem Beitrag wurde ein Gesamtsystem vorgestellt, welches in der Lage ist anhand von HD-Karten, einer hochgenauen Lokalisierung und Mehrfachbefahrungen skalierbar und automatisiert Trainingsdaten zu generieren. In Abbildung 3 ist in zwei Beispielen die Rückprojektion der Kartenmerkmale unter Berücksichtigung der statischen und dynamischen Verdeckung dargestellt. Abbildung 4 zeigt hingegen, wie durch die Multi-Drivekartierung die Menge der automatisch generierten Trainingsdaten je nach Bedarf skaliert werden kann. Die Qualität der Annotation hängt von verschiedenen Komponenten des Systems ab: (1) Genauigkeit der HD-Karte inklusive der 3D-Position der Kartenelemente, (2) Genauigkeit der 6D-Pose des Fahrzeugs in der HD Karte, (3) Genauigkeit der Kalibrierung der Sensoren und Rückprojektion in die Sensordomäne und (4) Genauigkeit der Verdeckungserkennung.

Um die Stärke des Ansatzes besser nutzen zu können, sollten Befahrungen mit einer großen Bandbreite von Verkehrssituationen, Wetter- und Lichtverhältnissen derselben Szenerie verwendet werden, um die Diversität des Datensatzes zu erhöhen. Aus der damit verbundenen zeitversetzten Aufnahme der Mehrfachbefahrungen ergeben sich Anforderungen an die Langzeitstabilität der Lokalisierungsmerkmale und die Möglichkeit der HD-Kartenverifikation.

Eine weitere Anwendung der Methode ist die Möglichkeit, semantisches Wissen aus einzelnen Sensordomänen automatisiert zu kartieren und durch das Lernen aus Karten das semantische Wissen auf das neue Sensor-Setup eines Fahrzeuges zu übertragen. Hierdurch könnte der Aufwand für die Erstellung von Trainingsdaten für neue Fahrzeugkonfigurationen stark reduziert werden.

Literatur

- [1] Esri *World Imagery*, https://\{switch:services,server\}arcgisonline.com/arcgis/rest/services/World_Imagery/MapServer/tile, © Esri, DigitalGlobe, GeoEye, i-cubed, USDA FSA, USGS, AEX, Getmapping, Aerogrid, IGN, IGP, swiss-topo, and the GIS User Community. Accessed 15. August 2023 via JOSM.
- [2] C. Cordts, M. Omran, S. Ramos, T. Rehfeld, M. Enzweiler, R. Benenson, U. Franke, S. Roth und B. Schiele, *The Cityscapes Dataset for Semantic Urban Scene Understanding*, International Conference on Computer Vision and Pattern Recognition (CVPR), 2016.
- [3] J. Xie, M. Kiefel, M. T. Sun und A. Geiger, *Semantic Instance Annotation of Street Scenes by 3D to 2D Label Transfer*, International Conference on Computer Vision and Pattern Recognition (CVPR), 2016.
- [4] H. Ling, S. Fidler, *Fast Interactive Object Annotation with Curve-GCN*, International Conference on Computer Vision (ICCV), 2019.
- [5] S. R. Richter, V. Vineet, S. Roth und V. Koltun, *Playing for Data: Ground Truth from Computer Games*, European Conference on Computer Vision (ECCV), 2016.

- [6] G. Ros, L. Sellart, J. Materzynska, D. Vazquez und A. M. Lopez, *The SYNTHIA Dataset: A Large Collection of Synthetic Images for Semantic Segmentation of Urban Scenes*, International Conference on Computer Vision and Pattern Recognition (CVPR), 2016.
- [7] A. Dosovitskiy, G. Ros, F. Codevilla, A. M. Lopez und V. Koltun, *CARLA: An Open Urban Driving Simulator*, CoRR, 2017.
- [8] Y. Liao, J. Xie und A. Geiger, *KITTI-360: A Novel Dataset and Benchmarks for Urban Scene Understanding in 2D and 3D*, Pattern Analysis and Machine Intelligence (PAMI), 2022.
- [9] B. Wilson, W. Qi, T. Agarwal, J. Lambert, J. Singh, S. Khandelwal, B. Pan, R. Kumar, A. Hartnett, J. Kaesemodel Pontes, D. Ramanan, P. Carr und J. Hays, *Argoverse 2: Next Generation Datasets for Self-Driving Perception and Forecasting*, arXiv, 2023.
- [10] H. Caesar, V. Bankiti, A. H. Lang, S. Vora, V. Liong, Q. Xu, A. Krishnan, Y. Pan, G. Baldan und O. Beijbom, *nuScenes: A multimodal dataset for autonomous driving*, International Conference on Computer Vision (ICCV), 2019.
- [11] TuSimple, *TuSimple benchmark*, <https://github.com/TuSimple/tusimple-benchmark>, accessed August 2023.
- [12] Ö. Ş. Taş, N. O. Salscheider, F. Poggenhans, S. Wirges, C. Bandera, M. R. Zofka, usw. *Making Bertha Cooperate-Team AnnieWAY's Entry to the 2016 Grand Cooperative Driving Challenge*, Transactions Intelligent Transportation Systems, 2018.
- [13] T. Strauß, J. Ziegler, J. Beck, *Calibrating Multiple Cameras with Non-Overlapping Views using Coded Checkerboard Targets*, Intelligent Transportation Systems Conference (ITSC), 2014.
- [14] J. Kümmerle, T. Kühner, M. Lauer, *Automatic Calibration of Multiple Cameras and Depth Sensors with a Spherical Target*, International Conference Intelligent Robots and Systems (IROS), 2018.
- [15] M. Sons, C. Stiller, *Efficient Multi-Drive Map Optimization towards Life-long Localization using Surround View*, International Conference on Intelligent Transportation Systems (ITSC), 2018.
- [16] F. Poggenhans, J.-H. Pauls, J. Janosovits, S. Orf, M. Naumann, F. Kuhnt und M. Mayr, *Lanelet2: A high-definition map framework for the future of automated driving*, International Conference on Intelligent Transportation Systems (ITSC), 2018.
- [17] A. Meyer, P. Skudlik, J.-H. Pauls und C. Stiller, *YOlinO: Generic Single Shot Polyline Detection in Real Time*, International Conference on Computer Vision Workshops, 2021.
- [18] L. Porzi, S. Rota Bulò, A. Colovic und P. Kotschieder, *Seamless Scene Segmentation*, International Conference on Computer Vision and Pattern Recognition (CVPR), 2019.

- [19] G. Neuhold, T. Ollmann, S. Rota Bulò und P. Kotschieder, *The Mapillary Vistas Dataset for Semantic Understanding of Street Scenes*, International Conference on Computer Vision (ICCV), 2017.

Random-Finite-Set-basiertes Multisensor-Multiobjekttracking für automatisierte und vernetzte Fahrzeuge: Aktueller Forschungsstand und offene Fragen

Martin Herrmann*, Charlotte Hermann*, Klaus Dietmayer[†]
und Michael Buchholz[‡]

Zusammenfassung: Die zeitliche und räumliche Verfolgung anderer Verkehrsteilnehmer ist ein essentieller Bestandteil automatisierter und vernetzter Fahrzeuge sowie vernetzter Infrastruktur-sensorsysteme und wird häufig mit Random-Finite-Set-basierten Methoden gelöst. Dabei stellt die Assoziation der beobachteten Objekte mit den Messungen insbesondere im Multisensorfall eine große algorithmische Herausforderung dar. Dieser Artikel fasst die wissenschaftlich beschriebenen Lösungsmethoden hierfür zusammen und ordnet insbesondere die in letzter Zeit an unserem Institut entwickelten parallelisierbaren Methoden auf Basis der Bayes-Parallel-Combination-Rule in den wissenschaftlichen Kontext ein. Letztere betrachten wir tiefergehend und diskutieren offene Fragestellungen und potenzielle Lösungsansätze.

Schlüsselwörter: Bayes-Parallel-Combination-Rule, Generalized-Labeled-Multi-Bernoulli, Labeled-Multi-Bernoulli, Product-Multi-Sensor-Multi-Object Tracking, Random-Finite-Sets

1 Einleitung

Jedem Autofahrer dürfte intuitiv bewusst sein, dass eine erfolgreiche Umgebungswahrnehmung zur Erfüllung der Fahraufgabe essentiell ist. Automatisierte Fahrzeuge sind für die Perception auf interne Sensorik angewiesen [1], die im Fall vernetzter Fahrzeuge durch externe Informationen anderer Fahrzeuge oder vernetzter Infrastruktur ergänzt werden kann. Die aus den Sensoren gewonnenen Daten werden in gängigen Architekturen unter anderem von Filteralgorithmen verarbeitet, die eine zeitliche und räumliche Zuordnung zwischen den Messungen und den beobachteten Verkehrsteilnehmern herstellen. Auf Basis der so gewonnenen Information über die vergangenen Trajektorien anderer Verkehrsteilnehmer kann ein automatisiertes Fahrzeug seine eigene Trajektorie planen und sogar Rückschlüsse auf das voraussichtliche zukünftige Verhalten der anderen ziehen.

*M. Herrmann und C. Hermann sind akademische Mitarbeitende am Institut für Mess-, Regel- und Mikrotechnik der Universität Ulm, Albert-Einstein-Allee 41, 89081 Ulm (martin.herrmann@uni-ulm.de; charlotte.herrmann@uni-ulm.de).

[†]K. Dietmayer ist Leiter des Instituts für Mess-, Regel- und Mikrotechnik der Universität Ulm, Albert-Einstein-Allee 41, 89081 Ulm (klaus.dietmayer@uni-ulm.de).

[‡]M. Buchholz ist akademischer Oberrat am Institut für Mess-, Regel- und Mikrotechnik der Universität Ulm, Albert-Einstein-Allee 41, 89081 Ulm (michael.buchholz@uni-ulm.de).

Um die Messungen aller Sensoren zu verarbeiten, kommen sogenannte Multiobjekttracker zum Einsatz. Während die Literatur mit den Nearest-Neighbor- (NN), den Probabilistic-Data-Association- (PDA) und den Multi-Hypothesis-Tracking-Verfahren (MHT-Verfahren) [2] eine Vielzahl unterschiedlicher Methoden bereitstellt, beschränkt sich dieser Artikel im Folgenden nur auf die Random-Finite-Set-basierten (RFS-basierten) Methoden. Genauer gesagt betrachten wir ausschließlich das Generalized-Labeled-Multi-Bernoulli-Filter (GLMB-Filter) [3] sowie dessen Approximation, das Labeled-Multi-Bernoulli-Filter (LMB-Filter) [4]. Beide nutzen Mahlers Finite-Set-Statistics-Framework (FISST-Framework) [5, 6] und können über die Multiobjekt-Bayesformel hergeleitet werden. Überdies handelt es sich beim GLMB-Filter um das bisher einzige rechentechnisch implementierbare Multiobjektfilter, für das Bayes-Optimalität eindeutig gezeigt wurde [7].

Trotz geeigneter Approximationsverfahren ist das GLMB-Filter in vielen, insbesondere komplexen Szenarien rechentechnisch herausfordernd. Sowohl in der Prädiktion als auch im Update entstehen eine Vielzahl unterschiedlicher Hypothesen, die in der Folge betrachtet werden müssen. Folglich wird häufig, so auch in unseren Anwendungen im automatisierten und vernetzten Fahren [8–10], das LMB-Filter genutzt, welches besonders die Prädiktion stark vereinfacht. Es hat unter anderem deshalb enorme rechentechnische Vorteile bei nur sehr geringen Leistungseinbußen. Unglücklicherweise stößt aber auch das LMB-Filter im Multisensorfall schnell an seine Grenzen, was sich auch in unseren Versuchen an der Infrastrukturanlage in Ulm-Lehr zeigt, bei der zusätzlich ein relativ komplexes Messmodell für Objektreferenzpunktmessungen zur Anwendung kommt [11]. Die Gründe hierfür liegen insbesondere in der fehlenden Möglichkeit, Messungen mehrerer Sensoren parallel zu verarbeiten, und in der NP-Schwere des Messung-zu-Track-Assoziationsproblems [12].

Dieser Beitrag widmet sich den Methoden, die verschiedene Forschungsgruppen zur Lösung dieser Herausforderungen in den letzten Jahren entwickelt haben. Wir konzentrieren uns dabei ausschließlich auf zentralisierte Ansätze, da die dezentralen Varianten unter Verwendung der suboptimalen Track-to-Track-Fusion ein deutlich geringeres Leistungsniveau zeigen [2, 13]. Anhand praktischer Anwendungsfälle erörtern wir noch offene Fragestellungen, die einer praxistauglichen Implementierung entgegenstehen, und präsentieren potenzielle Lösungsansätze.

2 Zentralisiertes Multisensor-Multiobjekttracking

Nach Bar-Shalom [2] unterscheidet man grob drei unterschiedliche Fusionsansätze: Messdatenfusion, zentralisiertes Multisensortracking und Track-to-Track-Fusion. Diese unterscheiden sich grundlegend anhand der fusionierten Datentypen. Bei der Messdatenfusion kombiniert ein vorgelagertes Modul die Sensordaten zu einer Art Supermessung, welche dann zeitlich gefiltert wird. Demgegenüber übernimmt der zentralisierte Multisensortracker die Fusion der Sensordaten implizit bei der Verarbeitung der Sensormessungen. Im letzten Fall steht allen Sensoren ein eigener Tracker zur Verfügung, deren gefilterte Daten anschließend fusioniert werden.

Wie bereits eingangs erwähnt, steht die Track-to-Track-Fusion den anderen Verfahren hinsichtlich der Leistungsfähigkeit stark nach, hat aber auch völlig andere Einsatzgebiete, beispielsweise in verteilten Architekturen mit wechselnden Netzwerktopologien und Kommunikationsteilnehmern. Die beiden anderen Methoden sind jedoch heutzutage nur noch schwer auseinanderzuhalten und oft in Mischformen präsent, weshalb generelle Aussagen

schwierig zu treffen sind. Jedoch sind Multisensorszenarien auf Basis reiner Messdatenfusion im betrachteten Anwendungsfeld automatisiertes und vernetztes Fahren eher untypisch, da sich beispielsweise Sensorsichtbereiche meist nicht vollständig überlappen oder die Modellierung für unterschiedliche Sensorprinzipien kompliziert und stark vom Sensortyp abhängig ist. Folglich hat man auch bei vorfusionierten Sensordaten oft einen Multisensortracker, der die Daten mehrerer Sensoren oder Sensorgruppen filtert und fusioniert.

In dieser Algorithmeklasse differenzieren wir erneut drei Methoden, nämlich solche mit Multisensor-Messmodell, serielle Ansätze und parallelisierbare Ansätze. Diese klassifizieren wir im Folgenden jeweils kurz und stellen die jeweiligen Vertreter mit entsprechender (Multiobjekt-)Wahrscheinlichkeitsdichtefunktion (WDF) [5] ausführlich vor, nachdem wir die grundlegende Funktionsweise der GLMB- und LMB-Filter zusammengefasst haben. Die Erläuterungen sind mit Absicht knapp und oberflächlich gehalten und nicht vollständig. Sie sollen auch fachlich weniger versierten Lesern einen groben Überblick erlauben und ein Grundverständnis für die Unterschiede der anschließend diskutierten Methoden vermitteln. Fachlich interessierte Leser verweisen wir auf [6] für tiefergehende Informationen.

2.1 Das GLMB- und das LMB-Filter

Eine GLMB-WDF $\pi(\mathbf{X})$ mit Zustands-RFS $\mathbf{X}$ hat im Allgemeinen die folgende Form [3]:

$$\pi(\mathbf{X}) = \Delta(\mathbf{X}) \sum_{c \in \mathbb{C}} w^{(c)}(\mathcal{L}(\mathbf{X})) [p^{(c)}(\cdot)]^{\mathbf{X}}. \quad (1)$$

Sehr vereinfacht gesagt modelliert diese alle möglichen Evolutionen eines Multiobjekt-zustands-RFS durch Hypothesen. Diese sind durch eine eindeutige Nummerierung gekennzeichnet und werden durch die Indexmenge $\mathbb{C}$ beschrieben. Die Summe in (1) bildet also diese möglichen Hypothesen ab, die, abhängig von der Realisierung, jeweils aus einer Eintrittswahrscheinlichkeit $w^{(c)}(\mathcal{L}(\mathbf{X}))$ und je einer WDF $p^{(c)}(\mathbf{X})$ je enthaltenem Objekt bestehen. Das Zustands-RFS $\mathbf{X}$ und dessen GLMB-WDF $\pi(\mathbf{X})$ sind fettgedruckt um hervorzuheben, dass es sich um eine indizierte (*labeled*) Größe handelt. Dabei wird jedem Objekt eine eindeutige Kennzeichnung zugeordnet und man spricht dann häufig auch von geschätzten Trajektorien anstelle eines geschätzten Zustands.

Im Updateschritt des GLMB-Filters müssen nun alle Hypothesen mit den Sensormessungen aktualisiert werden. Hierzu erzeugt man sämtliche mögliche Assoziationen zwischen den beobachteten Objekten $\mathbf{X}$ des RFS mit den einzelnen Messungen des Mess-RFS Z . Hierbei wächst die Anzahl der modellierten Hypothesen exponentiell an, da dieser Schritt für alle prädizierten Hypothesen (künftig *Quellhypothesen* genannt) durchgeführt werden muss und deren Anzahl weder in der Prädiktion noch im Updateschritt verringert werden kann. Aus diesem Grund greift man typischerweise zu einem Trick und berechnet nur die k wahrscheinlichsten Hypothesen, was man durch Lösung des k -rangiges Zuordnungsproblem (engl.: k -ranked assignment problem) erreicht. Dieses wird für alle relevanten Quellhypothesen separat berechnet, was man als Truncation bezeichnet und einer Approximation entspricht. Im zweidimensionalen Fall kann dieses Zuordnungsproblem in polynomieller Zeit durch Murty's Algorithmus gelöst werden (kubische Komplexität) oder mithilfe des Gibbs-Samplers (quadratische Komplexität) angenähert werden [14]. Im mehrdimensionalen Fall dagegen ist die Lösung NP-schwer und ist damit im Allgemeinen nicht in endlicher Zeit berechenbar [12]. Abschließend muss jedoch in jedem Fall die Zahl der Hypothesen

insgesamt beschränkt werden, was man als Pruning bezeichnet. Hierfür sortiert man alle aktualisierten Hypothesen nach ihrem Gewicht und entfernt die unwahrscheinlichen.

Das LMB-Filter unterscheidet sich nur geringfügig davon, denn eine LMB-WDF stellt einen Spezialfall einer GLMB-WDF mit nur einer Komponente dar. Beschreibt nun also eine LMB-WDF den Multiobjektzustand zu einem Zeitpunkt k , für welche dann das Bayes-Update durchgeführt wird, so ergibt sich im Ergebnis eine GLMB-WDF. Damit wird klar, dass der Updateschritt des LMB-Filters prinzipiell dem des GLMB-Filters gleicht. Allerdings approximiert man anschließend die GLMB-WDF wieder als LMB-WDF. Kurz gesagt verliert man in diesem Schritt modellierte Abhängigkeiten zwischen den Objekten, behält aber die gesamte räumliche Information über die Objektzustände bei, indem man sie in eine Mixturdichte komprimiert. Im Ergebnis reduziert sich die Komplexität der WDF deutlich, da die Zahl der Hypothesen enorm abnimmt, weshalb das LMB-Filter im Allgemeinen deutlich recheneffizienter ist. Zusätzlich vereinfacht sich der Prädiktionsschritt, da die Objekte als unabhängig betrachtet werden.

2.2 RFS-basiertes Multisensor-Multiobjekttracking

Abbildung 1 zeigt die Strukturbilder der im Folgenden vorgestellten Methoden. Die Farben markieren den Typ der WDF [5], der in den jeweiligen Teilen genutzt wird. Dabei steht rot für GLMB- und blau für LMB-WDFs.

2.2.1 Ansätze mit Multisensor-Messmodell

Im allgemeinen Fall kann das Multisensor-Multiobjektupdate mit RFS-basierten Methoden mit der Multisensor-Bayesgleichung [5] berechnet werden:

$$\pi(\mathbf{X}|Z) \propto \pi(\mathbf{X}) \cdot g(Z|\mathbf{X}). \quad (2)$$

Dabei bezeichnet $g(Z|\mathbf{X})$ das Multisensor-Multiobjektmessmodell, welches zumindest theoretisch auch für stark korrelierte Multisensormessungen immer gebildet werden kann [2]. Jedoch führt dies unausweichlich auf ein NP-schweres, mehrdimensionales Zuordnungsproblem ohne gemeingültige, technisch realisierbare Lösungsmethode [12]. Theoretisch betrachtet lassen sich sowohl das GLMB- als auch das LMB-Filter mit einem Multisensor-Messmodell darstellen, wie in den Abbildungen 1a und 1b dargestellt. In der Praxis sind nur Implementierungen mit stark begrenzter Sensoranzahl V in bestimmten Szenarien möglich, beispielsweise in nicht echtzeitkritischen Anwendungen wie der Nachbearbeitung von Datensätzen.

Das **GLMB-Filter mit Multisensor-Messmodell** wurde erstmals für zwei Sensoren in [15] präsentiert, kann aber theoretisch leicht um beliebige Sensoren erweitert werden [12, 17]. Zumindest theoretisch ist es das leistungsstärkste Filter, verliert diese Eigenschaft aber aufgrund der notwendigen Approximationen in der Praxis oft. Nicht wirklich anders stellt sich die Situation für das **LMB-Filter mit Multisensor-Messmodell** dar. Es ist, vermutlich deshalb, auch in der Literatur nicht beschrieben, aber in logischer Konsequenz aus dem LMB-Filter und dem GLMB-Filter mit Multisensor-Messmodell ableitbar. Anders ausgedrückt berechnet das LMB-Filter mit Multisensor-Messmodell die Prädiktion auf der LMB-WDF und nutzt für das Multisensorupdate das Update des GLMB-Filters mit Multisensor-Messmodell. Das Multisensorupdate muss also wie

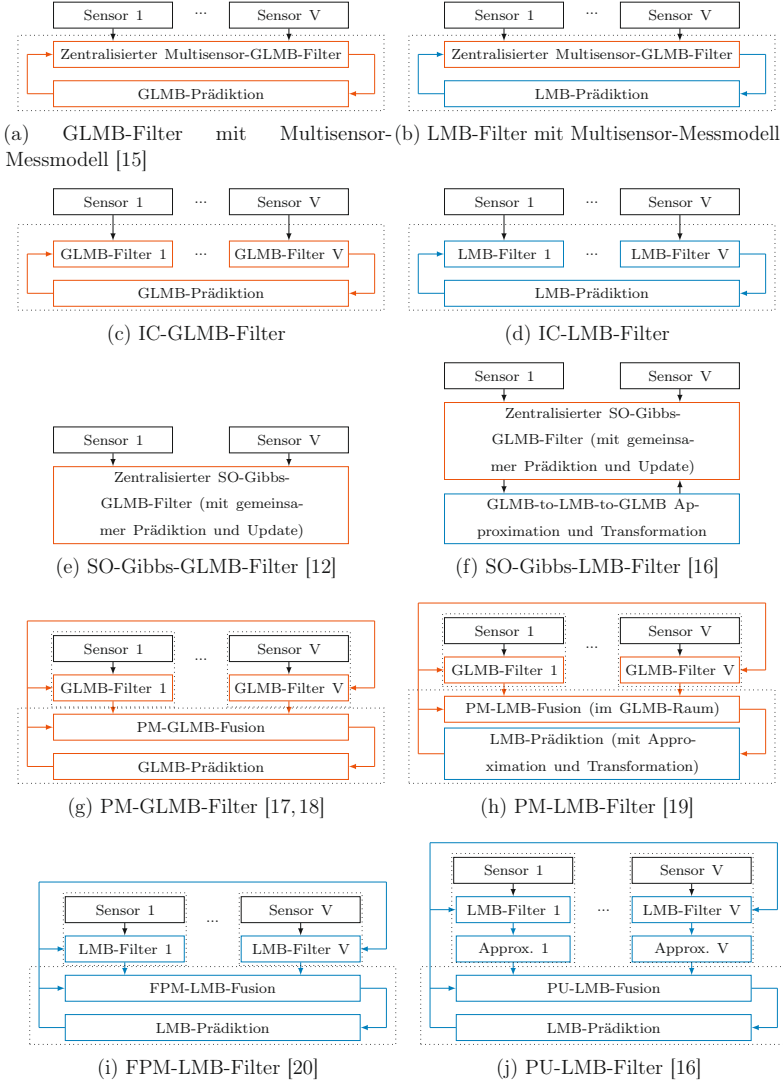


Abbildung 1: Übersicht der zentralisierten Multisensor-Multiobjekttracking-Ansätze, wobei die Farben den Typ der verwendeten WDF auf der Schnittstelle oder innerhalb der Module indizieren: Rot steht für GLMB-WDFs und blau für LMB-WDFs.

üblich im GLMB-Raum durchgeführt werden. Auch dieser Filteransatz liefert zumindest theoretisch die bestmögliche Güte, ist aber ebenso nur in speziellen Fällen überhaupt anwendbar. Zusätzlich leiden beide Ansätze unter dem Problem, dass Lösungsmethoden für das mehrdimensionale Zuordnungsproblem nicht sinnvoll parallelisierbar sind.

2.2.2 Serielle Ansätze

Serielle Ansätze können angewendet werden, wenn sich das Multisensor-Messmodell $g(Z|\mathbf{X})$ als Produkt der einzelnen Sensormessmodelle schreiben lässt, also die Messungen unabhängig voneinander sind [12]. Diese Annahme kann in der Praxis häufig getroffen werden, da die Unabhängigkeit der Messungen einzig vom Messrauschen der Sensoren abhängt. Dann gilt für das Bayes-Update

$$\pi(\mathbf{X}|Z) \propto \pi(\mathbf{X}) \cdot \prod_{s=1}^V g^{(s)}(Z^{(s)}|\mathbf{X}), \quad (3)$$

wodurch die Klasse der Iterated-Corrector-Ansätze (IC-Ansätze) sowie der Multisensor-Gibbs-Sampler begründet wird.

Abbildung 1c und 1d zeigen die Filterstrukturen der **IC-Ansätze Iterated-Corrector-GLMB-Filter (IC-GLMB-Filter)** und **Iterated-Corrector-LMB-Filter (IC-LMB-Filter)**. Beide ähneln sich in ihrer Struktur insofern, als dass alle Sensormessungen sequentiell verarbeitet werden. Das IC-GLMB-Filter ist theoretisch betrachtet äquivalent zum GLMB-Filter mit Multisensor-Messmodell und damit optimal, wohingegen das IC-LMB-Filter das LMB-Filter mit Multisensor-Messmodell approximiert, indem die GLMB-zu-LMB-Approximation nach jedem Sensorupdate durchgeführt wird. Aber auch im IC-GLMB-Filter müssen die gängigen Approximationen nach jedem Update durchgeführt werden, da die Zahl der Hypothesen ansonsten zu stark ansteigen würde. Damit erfordern beide Ansätze maximal Lösungsalgorithmen für das zweidimensionale Zuordnungsproblem und sind technisch realisierbar. Allerdings wird die Filterleistung abhängig von der Verarbeitungsreihenfolge der Sensormessungen und sinkt im Allgemeinen erkennbar [12]. Zudem bleibt der Nachteil der fehlenden Parallelisierbarkeit der Messdatenverarbeitung bestehen.

Abgesehen von den parallelisierbaren Methoden, stellen das IC-GLMB- und insbesondere das IC-LMB-Filter die Implementierungen mit den kürzesten Rechenzeiten dar. Dies gilt umso mehr, wenn zur Lösung des Zuordnungsproblems Sampling-basierte Methoden, wie z.B. der Gibbs-Sampler mit quadratischer [14, 21] bzw. linearer [22] Komplexität, verwendet werden. Zudem ist die Umsetzung der IC-Varianten, beispielsweise auf Basis von Implementierung für Einzelsensoren, sehr einfach. Schließlich erlauben diese Ansätze auch die Anwendung von *Grouping*-Strategien, also der unabhängigen und parallelen Berechnung der Messupdates unabhängiger Objektgruppen [4, 23].

Auswertungen in [12, 17–20] zeigen aber eindrücklich die Schwierigkeiten der iterativen Ansätze, die vor allem bei ungünstiger Verarbeitungsreihenfolge kritische Folgen haben können. Deutlich besser schlägt sich das **Sub-optimal Gibbs-Sampling-based GLMB-Filter (SO-Gibbs-GLMB-Filter)** [12], dessen Filterstruktur in Abbildung 1e abgebildet ist. Trotz eines Multisensorupdates zählen wir es zu den seriellen Ansätzen, da es auf derselben Annahme unabhängiger Sensormessungen beruht. Zusätzlich verwendet dieses Filter einen kombinierten Prädiktions- und Updateschritt (engl. *joint-prediction-and-update*), der das Filter sehr leistungsfähig und effizient hinsichtlich der benötigten Hypothesen zur Abbildung der mehrdeutigen Messassoziation macht.

Gegenüber den iterativen Varianten benötigt das SO-Gibbs-GLMB-Filter allerdings eine deutlich höhere Rechenzeit, was insbesondere für die ebenfalls bekannte optimale Variante, das **Optimal Gibbs-Sampling-based GLMB-Filter (O-Gibbs-GLMB-Filter)** gilt, das aber ebenfalls in die Kategorie der NP-schweren Probleme fällt [12] und deshalb hier nur am Rand erwähnt wird. Es ist bisher auch nicht bekannt, inwiefern Grouping-Strategien mit dem Multisensor-Gibbs-Ansatz kombiniert werden können oder ob eine Multisensorvariante des Gibbs-Samplers mit linearer Komplexität möglich ist. Beide Erweiterungen dürften bei erfolgreicher Realisierung die Bedeutung des Ansatzes aufgrund der zu erwartenden Beschleunigung der Rechenzeit deutlich erhöhen. Bekannt dagegen ist die entsprechende Anwendung auf LMB-WDFs, also das **Sub-optimal Gibbs-Sampling-based LMB-Filter (SO-Gibbs-LMB-Filter)**, dessen Struktur in Abbildung 1f dargestellt ist [16].

2.2.3 Parallelisierbare Ansätze

Ebenfalls unter der Bedingung unabhängiger Sensormessungen kann die Bayes-Updategleichung geschickt so umgestellt werden, sodass ein paralleles und unabhängiges Update der Sensormessungen möglich wird. Diese Formulierung ist bekannt als Bayes-Parallel-Combination-Rule (BPCR) [24] und lautet für RFSs

$$\pi(\mathbf{X}|Z) \propto \left(\pi(\mathbf{X})\right)^{1-V} \prod_{s=1}^V \pi^{(s)}(\mathbf{X}|Z^{(s)}). \quad (4)$$

Auf ihr basiert beispielsweise auch das Information-Matrix-Fusion-Filter (IMF-Filter) [2], weshalb wir die folgenden Filter auch als Multisensor-Informationsfilter oder Multisensor-IMF-Äquivalent bezeichnen. Ganz generell zeichnet sich die BPCR durch die Parallelisierbarkeit der Sensorupdates und damit durch eine mögliche verteilte Implementierung aus. Damit ist auch ein Intellectual-Property-Schutz (IP-Schutz) der Messmodelle möglich. Allerdings kann die BPCR nur unter zwei Bedingungen angewendet werden: Es wird erstens eine globale Fusionsinstanz benötigt, die allen beteiligten Sensoren eine global gültige Prädiktion bereitstellt. Und zweitens muss das Produkt der lokal aktualisierten WDFs durch das $(V - 1)$ -fache Produkt der prädierten WDF teilbar sein, was die Anwendung der BPCR auf GLMB- und LMB-WDFs herausfordernd macht. Dennoch konnten sowohl die Forschungsgruppe um Robertson [16] als auch wir in den vergangenen Jahren entsprechende Lösungen für beide Dichten finden.

Den ältesten der vier Ansätze stellt das in Abbildung 1g skizzierte **Product-Multi-Sensor-GLMB-Filter (PM-GLMB-Filter)** [17, 18] dar. Bei diesem berechnen die lokalen Filter ihr Sensorupdate jeweils im GLMB-Raum auf Basis der prädierten globalen GLMB-WDF. Die Filterformulierung kann grundsätzlich exakt hergeleitet werden, weshalb das Filter neben dem GLMB-Filter mit Multisensor-Messmodell die einzige Bayes-optimale Multisensor-Implementierung darstellt. Jedoch ignoriert man in der Praxis sinnvollerweise die Kreuzkovarianzen der einzelnen Sensorinnovationen, da deren Berechnung numerisch instabil ist [18].

Dennoch zeigt das Filter in Simulationen sehr gute Leistungen (wobei ein direkter Vergleich bisher nur mit dem IC-GLMB- und dem IC-LMB-Filter durchgeführt wurde) und stellt aufgrund der Parallelisierbarkeit die schnellste bekannte Multisensor-GLMB-Filtervariante dar. Nachteilig ist aber die in der Praxis recht komplizierte Parametrisierung der Truncation- und Pruningparameter, die deutliche Auswirkungen auf die Filterleistung

haben. Insbesondere zeigt das PM-GLMB-Filter Schwierigkeiten mit der Geburt neuer Objekte, wenn die zentrale Instanz ein Pruning auf die prädierte WDF anwendet. Im Gegenzug zu den vorigen Filtern ist dies aber in der Praxis oft notwendig.

Das **Product-Multi-Sensor-LMB-Filter (PM-LMB-Filter)** [19], dessen Strukturdiagramm in Abbildung 1h dargestellt ist, bildet eine mögliche Alternative auf Basis der LMB-WDFs ab. In diesem Fall wird die Prädiktion vereinfacht im LMB-Raum durchgeführt, aber die Messupdates und die Filterfusion finden im GLMB-Raum statt. Dies entspricht mathematisch betrachtet dem LMB-Filter mit Multisensor-Messmodell, und zeigt dementsprechend gute Ergebnisse in der Praxis. Außerdem ist das PM-LMB-Filter noch einmal deutlich schneller als das PM-GLMB-Filter. Einschränkung muss allerdings erwähnt werden, dass die Rechenzeiten des IC-LMB-Filters nur dann erreicht werden, wenn dieser kein Grouping verwendet, da ein Grouping für die parallelisierbaren Methoden noch nicht entwickelt wurde, wenngleich es zumindest für die LMB-Varianten möglich erscheint.

Im Gegensatz zu den beiden vorigen Varianten führt das **Parallel-Update LMB-Filter (PU-LMB-Filter)** [16] den Filterzyklus vollständig im LMB-Raum durch (wobei wir die kurzzeitige Konversion in den GLMB-Raum für das Einzelsensorupdate hier der Einfachheit halber vernachlässigen), wie in Abbildung 1j dargestellt. Damit müssen anstelle der Hypothesen nur die jeweiligen Objekte mit demselben Label fusioniert werden, deren Anzahl üblicherweise deutlich geringer ist als die der Hypothesen bei der Fusion im GLMB-Raum. Allerdings müssen die lokal aktualisierten räumlichen WDFs der Objekte vor ihrer Fusion approximiert werden um die Division durch die prädierte WDF durchführen zu können. Dabei müssen aus den normalverteilten Mixturdichten Normalverteilungen geschätzt werden. Im einfachsten Fall nimmt man dazu die wahrscheinlichste Mixturkomponente oder das gewichtete Mittel, die Autoren stellen aber auch komplexere Methoden zur Approximation zur Verfügung.

Das neueste und schnellste Filter, ist das **Fast-Product-Multi-Sensor-LMB-Filter (FPM-LMB-Filter)** [20], dessen Struktur in Abbildung 1i dargestellt ist. Es baut direkt auf der Filterformulierung des PU-LMB-Filters auf, löst aber das Divisionsproblem, indem zur Fusion der räumlichen Dichten auf das IMF zurückgegriffen wird. So bleibt auch das FPM-LMB-Filter vollständig im LMB-Raum und ist deshalb wie das PU-LMB-Filter noch einmal deutlich weniger rechenintensiv als die anderen parallelisierbaren Methoden.

2.3 Diskussion

In der Praxis spielen die Ansätze mit Multisensor-Messmodell aufgrund ihrer rechen-technischen Komplexität eigentlich keine Rolle. Die seriellen Verfahren dagegen sind in Multisensor-Anwendungen quasi Standard und profitieren von ihrer programmatischen Einfachheit sowie ihrer relativ guten Geschwindigkeit und ihrer oft ausreichenden Leistungsfähigkeit.

Werden jedoch hohe Echtzeit- oder Leistungsanforderungen gestellt, kommen die alternativen Ansätze ins Spiel, und hier stellen die Product-Multi-Sensor-Ansätze (PM-Ansätze) einen guten Kompromiss dar. Ihre Eigenschaften erhalten sie zu großen Teilen aus ihrem differenzierten Herrsche-und-Teile-Ansatz, durch den sie sich von den anderen Varianten unterscheiden. Einerseits ermöglicht der Ansatz eine echte Multisensorfusion, die gleichzeitig alle Sensormessungen berücksichtigt und dabei enorm schnell ist. Dies erkauft man sich allerdings dadurch, dass die Truncation der lokalen Updates isoliert und

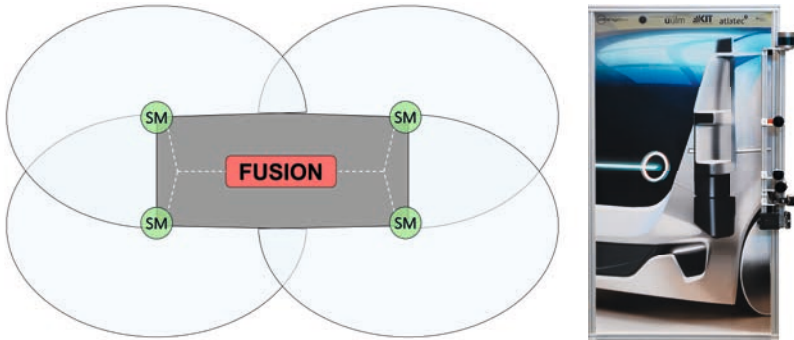


Abbildung 2: Perzeptionskonzept (links) und SM (rechts) aus UNICARagil

unabhängig durchgeführt wird. Allerdings sieht die Situation bei den IC-Filtern sehr ähnlich aus, die bei der Truncation ebenfalls nur begrenzte Informationen zur Verfügung haben, genauso wenig auf ein Multisensorupdate bauen können und nach jedem Sensorupdate ein Pruning benötigen. Einzig das SO-Gibbs-GLMB- und das SO-Gibbs-LMB-Filter haben hier Vorteile, die sich jedoch in einer deutlich höheren Rechenzeit niederschlagen. Andererseits ermöglichen die parallelisierbaren Filter als einzige eine verteilte Implementierung und den Schutz geistigen Eigentums der Sensorhersteller an deren Messmodellen.

Zusammengefasst lässt sich also sagen, dass die PM-Ansätze enorme Vorteile bieten und dabei kaum praktisch relevante Nachteile haben.

3 Offene Forschungsfragen und zukünftige Arbeiten

Gegenüber den anderen Methoden sind die parallelisierbaren Ansätze noch relativ jung und entsprechend wenig verbreitet. Einer solchen Verbreitung stehen neben fehlender Bekanntheit allerdings auch noch einige offene Forschungsfragen im Weg, die bisher noch nicht abschließend gelöst werden konnten. Sie sollen im Folgenden anhand zweier Multisensor-Anwendungsfälle aus unseren Forschungsprojekten erläutert und potenzielle Lösungsstrategien aufgezeigt werden.

3.1 Modulares Sensormodulkonzept aus UNICARagil

Modularität und Redundanz haben im Fahrzeugkonzept der im Rahmen des Projekts UNICARagil entwickelten Fahrzeuge einen großen Stellenwert [10]. Auch die Perception folgt diesen Projektzielen und besteht, wie in Abbildung 2 dargestellt, aus vier unabhängigen Sensormodulen (SMs) und einer zentralen Fusionseinheit. Es handelt sich damit um eine zentralisierte Netzwerktopologie, für welche die PM-Filter ideal geeignet wären. In den realisierten Fahrzeugen kommt allerdings eine suboptimale Track-to-Track-Fusion zum Einsatz, da einige offene Forschungsfragen bisher nicht beantwortet wurden.

Da die Sensoren auf einem beweglichen Fahrzeug montiert sind, benötigt man im Multiobjekttracking ein **adaptives Geburtenmodell**. Leider ist das gängige Modell für LMB- und GLMB-Filter sensorspezifisch und hauptsächlich für die IC-Ansätze geeignet [4].

Bei den PM-Ansätzen dagegen generieren alle Sensoren gleichzeitig mögliche Geburtskandidaten, wobei sich mit wachsender Sensoranzahl mit an absoluter Sicherheit grenzender Wahrscheinlichkeit einige davon doppelten. Um die Filtergüte nicht zu sehr zu schmälern benötigt es geeignete Methoden zur Vermeidung solcher Mehrfachgeburten. Eine potentiell vielversprechende Option wäre die Verwendung LMB-basierter Track-to-Track Methoden, welche diese Geburtskandidaten kombinieren. Im besten Fall würde man dann nicht nur die Mehrfachgeburt verhindern, sondern sogar von der Multisensorinformation profitieren.

Zusätzlich stellen aber auch die verbauten Radarsensoren eine Herausforderung dar. Prinzipbedingt sind sie nicht exakt zeitlich synchronisierbar, im Gegensatz zu den Kameras und Lidarsensoren. Um entsprechende Sensordaten fusionieren zu können, wäre eine Möglichkeit, die Sensordaten auf die synchronisierten Zeitpunkte des Filter-Updates zu präzisieren. Allerdings ist dies aufgrund der Unsicherheiten in den Messdaten sowie deren Unvollständigkeit in der Praxis oft schwierig. Alternativ könnte eine Fusion sogenannter Out-of-Sequence-Messungen realisiert werden. Hierzu existieren bisher jedoch noch keine Lösungen.

Letztlich leiden außerdem sowohl das PM-GLMB- als auch das PM-LMB-Filter in der Praxis unter der Vielzahl an zu berücksichtigenden Hypothesen. Beispielsweise führt die Vernachlässigung einer Quelhypothese durch nur einen einzigen Sensor in der Fusion zur vollständigen Auslöschung aller darauf aufbauenden Hypothesen. Gerade mit steigender Objekt- und Sensoranzahl nimmt die Wahrscheinlichkeit solcher Auslöschungen prinzipiell zu, und es braucht Methoden zur Steuerung. In unseren Implementierungen erreichen wir dies aktuell nur unzureichend durch eine Kombination mehrerer Strategien: Einerseits parametrisieren wir die Truncation der Sensoren so, dass deutlich mehr Hypothesen in den lokalen Updates berechnet werden als Hypothesen das Pruning nach der Fusion überleben. Andererseits beschränken wir oft auch die Anzahl der Hypothesen nach der Prädiktion durch einen zusätzlichen Pruningschritt. Und letztlich zwingen wir die lokalen Filter dazu, einen bestimmten Anteil ihrer verfügbaren Rechenzeit so auf die Quelhypothesen zu verteilen, dass möglichst alle mindestens einmal berücksichtigt werden. Dieser Ansatz ist natürlich rechentechnisch und leistungsmäßig fraglich und sollte verbessert werden, beispielsweise durch eine **verbesserte Koordination der zu berücksichtigenden Hypothesen, (also einer koordinierten Truncation)**, in den lokalen Filtern.

3.2 Multiobjekttracking für verteilte Infrastruktursensorsysteme

Insbesondere im unübersichtlichen und unsignalisierten innerstädtischen Verkehr tun sich automatisierte Fahrzeuge derzeit noch sehr schwer [9]. Im Projekt Lokales Umfeldmodell für das kooperative, automatisierte Fahren in komplexen Verkehrssituationen (LUKAS) nutzen wir ein Infrastruktursensorsystem zur Unterstützung vernetzter Fahrzeuge in unserer Versuchsanlage in Ulm-Lehr. Abbildung 3 zeigt auf der linken Seite eine Karte der Kreuzung und auf der rechten Seite zwei der exponiert angebrachten Kamerasensoren in der nördlichen Kreuzungzufahrt. Eine Besonderheit stellt dabei das verwendete Referenzpunktmessmodell zur Vermeidung systematischer Messfehler dar [11], bei dem Sensoren flexibel diejenigen Messdaten übermitteln können die auch real inferierbar sind. Typischerweise ist dies die Position eines Fahrzeugeckpunkts (sogenannter Referenzpunkt) und gegebenenfalls noch die Objektausdehnung in eine oder zwei Richtungen. Insgesamt besteht das System aus mehreren Kamera-, Radar- und Lidarsensoren, deren Messungen drahtlos an einen Multi-access-Edge-Computing-Server (MEC-Server) gesendet und dort fusioniert werden. Folglich



Abbildung 3: Simuliertes Verkehrsszenario an der Infrastrukturanlage in Ulm-Lehr mit sieben Fahrzeugen (farbige Linienzüge), 14 Sensoren (deren Sichtbereich blau dargestellt sind), allen Sensormasten (markiert durch $\blacktriangle$), den statischen Geburtsorten (markiert durch $\circ$) und dem ausgewerteten Bereich (rot gestrichelter Bereich) [25,26].

besitzt auch dieses System eine zentralisierte Netzwerktopologie, die sich hervorragend für die PM-Filter eignet.

Im Gegensatz zum UNICAR^{agil}-Fahrzeug ist deren Anwendung mit statischem Geburtenmodell in einer Digital-Twin-Simulation [27] hier aufgrund der statischen Sensorpositionierung möglich. Im gezeigten Szenario simulieren wir ausschließlich mehrere Kamerasensoren an den mit $\blacktriangle$ markierten Montagepositionen. Die jeweiligen Sensorsichtbereiche sind blau hervorgehoben und die rot umrandete Zone markiert den ausgewerteten Bereich. In der Simulation beschränken sich die Sensoren auf die Detektion der Eckpunkte, die sie mit einer normalverteilten und unabhängigen Messunsicherheit σ generieren. Damit fällt dem Multiobjektracker die Aufgabe der Ausdehnungsschätzung der Objekte zu, was möglich ist, solange mehrere Sensoren dasselbe Objekt an verschiedenen Referenzpunkten beobachten. Zudem generieren die Sensoren durchschnittlich einen zusätzlichen Falschalarm pro Zeitschritt und übersehen Fahrzeuge mit einer Wahrscheinlichkeit von 10 %.

Dieses Szenario haben wir in einer Monte-Carlo-Simulation mit 100 Wiederholungen ausgeführt und die Ergebnisse der vier Filtervarianten IC-GLMB-, IC-LMB-, PM-GLMB- und PM-LMB-Filter mit der OSPA⁽²⁾-Metrik [28] ausgewertet. Die Ergebnisse sind im oberen Diagramm der Abbildung 4 für die Messunsicherheiten $\sigma^2 = 0.5\text{m}^2$ und $\sigma^2 = 2\text{m}^2$ dargestellt. Im unteren Diagramm sieht man die Kardinalitätsschätzung für den zweiten Fall. Hierbei zeigen sich einige Besonderheiten. Entgegen der gängigen Annahme schlagen sich die LMB-Filter besser als die GLMB-Filter. Dies liegt daran, dass die Anzahl der eigentlich benötigten Hypothesen in den GLMB-Filtern derart hoch ist, dass eine Ausführung rechentechnisch nicht mehr realisierbar und schon gar nicht echtzeitfähig ist. Daher muss auf eine geringere Anzahl von Hypothesen reduziert werden. Konsequenterweise steigt mit steigender Messunsicherheit dann auch der Leistungsunterschied zwischen den beiden Filtertypen, weil die Messung-zu-Track-Assoziation mehrdeutiger wird.

Betrachtet man aber jeweils die Gruppe der GLMB- und LMB-Filter unabhängig voneinander fällt auf, dass die PM-Variante immer besser als die jeweilige IC-Variante abschneidet. Dies verdanken die PM-Filter ihrer überlegenen Kardinalitätsschätzung auf-

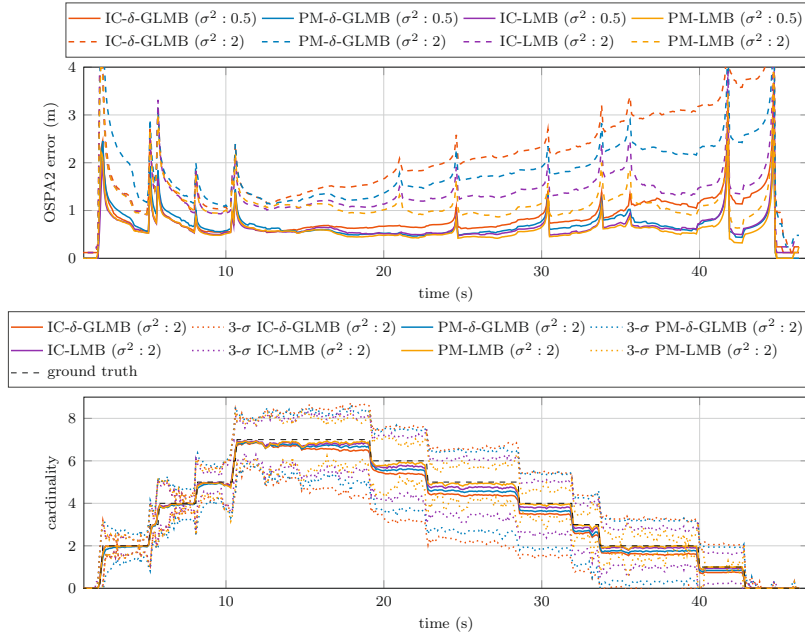


Abbildung 4: Optimal-Sub-pattern-Assignment⁽²⁾ (OSPA⁽²⁾)-Fehler aller Filter für beide simulierten Messunsicherheiten (oben), sowie die geschätzte Kardinalität der Filter für den Fall mit höherem Messrauschen (unten), jeweils gemittelt über 100 Simulationsdurchläufe.

grund des echten Multisensorupdates, wie man im unteren Diagramm deutlich sehen kann. Im Verlauf können die PM-Filter die Kardinalität immer etwas präziser schätzen, wobei die großen Fehler natürlich daher rühren, dass aufgrund des statischen Geburtenmodells einmal verlorengegangene Tracks dauerhaft verloren sind.

Anhand dieses Beispiels und der Ergebnisse lassen sich aber weitere noch offene Forschungsfragen ableiten, die in naher Zukunft beantwortet werden sollten. So ist beispielsweise das Multisensor-Update des PM-Filters aufgrund der idealen Parallelisierbarkeit sehr schnell berechenbar, aber in diesem Beispielszenario dennoch nicht ganz echtzeitfähig. Der Grund: Für die IC-Filter gibt es sogenannte **Groupingansätze**, wobei anhand der Unsicherheiten der Sensormessungen die WDF in unabhängige Teil-WDFs aufgeteilt wird. Dies führt, fast verlustfrei, zu einer enormen Verkürzung der Rechenzeit, da die Größe des Assoziationsproblems deutlich reduziert wird. Eine Erweiterung auf den Multisensor-Fall wäre folglich sehr attraktiv und könnte zumindest für das FPM-LMB- als auch das PU-LMB-Filter aufgrund der modellierten Objektunabhängigkeit möglich sein.

Anhand entsprechender effizienter Implementierungen wäre abschließend vor allem ein **fairer Vergleich der unterschiedlichen PM-Filteransätze sowohl untereinander, als auch gegenüber den Multisensor-Gibbs-Filtern** interessant. Denn es ist bisher nicht bekannt wie die Ansätze gegeneinander hinsichtlich verschiedener Kriterien, wie die Schätzungsgüte, Rechenzeit, etc., abschneiden. Hierfür wären allerdings vergleichbare und

rechentechnisch effiziente Implementierungen notwendig, die es leider bisher weder in der Open-Source-Szene, noch bei uns am Institut gibt.

Aufgrund der **numerisch instabilen Berechnungsvorschrift für die Innovationskovarianz in den PM-Filtern** müssen diese approximiert werden [18]. Praktisch vernachlässigt man diese Werte, da sie nachweislich kaum Relevanz haben. Dennoch könnte ihre Berechnung in spezifischen Situationen von Vorteil sein, weshalb die Entwicklung einer geschlossenen Berechnungsvorschrift vorteilhaft wäre.

Mit den Poisson Multi-Bernoulli Mixture (PMBM)-Filtervarianten für Trajektorien (siehe [29] für einen Überblick) existieren noch weitere vielversprechende RFS-basierte Filtertypen in der einschlägigen Literatur. Diese profitieren wie das LMB-Filter von einer vereinfachten Prädiktion, können aber dennoch Objektabhängigkeiten modellieren und werden in praktischen Untersuchungen oft als dem GLMB-Filter überlegen dargestellt [29]. Eine Anwendung der PMBM-WDF auf die BPCR und die entsprechende Entwicklung des **Product-Multi-Sensor-Poisson-Multi-Bernoulli-Mixture-Filters (PM-PMBM-Filters)** erscheint also vielversprechend.

Durch die PM-Filter gibt es nun erstmals die Möglichkeit, die Update-Berechnung im Multisensor-Multiobjekt-Fall parallel durchzuführen. Möchte man diese jedoch in einem Netzwerk verteilt berechnen lassen, müssen die Verbindungen breitbandig genug sein, um die jeweiligen WDFs zu übertragen. Hier stößt man sowohl bei GLMB- als auch bei LMB-WDFs derzeit noch schnell an praktische Grenzen, insbesondere in drahtlosen Netzwerken. Entsprechend müssen **geeignete Komprimierungsverfahren für WDFs** gefunden werden und der jeweilige Einfluss auf die Filterleistung untersucht werden.

4 Fazit

In diesem Artikel haben wir die aktuell wissenschaftlich beschriebenen RFS-basierten Multisensoransätze für GLMB- und LMB-WDFs vorgestellt und anhand ihrer Vor- und Nachteile verglichen. Im Ergebnis zeigt sich, dass die diskutierten PM-Filter aufgrund ihrer Laufzeit und Leistungsfähigkeit eine ernstzunehmende Alternative für die gängigen IC-, aber auch die Multisensor-Gibbs-Varianten darstellen. Allerdings stehen mit den Themen wie Grouping und adaptive Geburt noch gewichtige offene Fragen im Raum, deren Beantwortung wir in der kommenden Zeit angehen und auch andere Wissenschaftler dazu motivieren möchten. Denn im Multiobjektfall treten die Vorteile des parallelen Multisensor-Updates umso mehr zu Tage, die insbesondere in Anwendungen wie dem automatisierten Fahren mit vielen zu beobachtenden Verkehrsteilnehmern zum Tragen kommen.

Danksagung

Diese Arbeit wurde finanziell vom Bundesministerium für Wirtschaft und Klimaschutz (BMWK) im Rahmen des Programms „Hoch- und vollautomatisiertes Fahren in anspruchsvollen Fahrsituationen“ (Projekt LUKAS, FKZ 19A16010I) und vom Bundesministerium für Bildung und Forschung (BMBF) (Projekt UNICAR*agil*, FKZ 16EMO0290) unterstützt.

Literatur

- [1] H. Winner, S. Hakuli, F. Lotz, and C. Singer, Eds., *Handbuch Fahrerassistenzsysteme*. Springer Fachmedien, Wiesbaden, 2015.
- [2] Y. Bar-Shalom, P. K. Willett, and X. Tian, *Tracking and Data Fusion: A Handbook of Algorithms*. Connecticut: YBS Publishing, 2011.
- [3] B. T. Vo and B. N. Vo, “Labeled random finite sets and multi-object conjugate priors,” *IEEE Trans. Signal Process.*, vol. 61, no. 13, pp. 3460–3475, 2013.
- [4] S. Reuter, B. T. Vo, B. N. Vo, and K. Dietmayer, “The Labeled Multi-Bernoulli Filter,” *IEEE Trans. Signal Process.*, vol. 62, no. 12, pp. 3246–3260, 2014.
- [5] R. P. S. Mahler, *Statistical Multisource-Multitarget Information Fusion*. Artech House, Norwood, 2007.
- [6] —, *Advances in Statistical Multisource-Multitarget Information Fusion*. Artech House, Norwood, 2014.
- [7] —, “Exact closed-form multitarget bayes filters,” *Sensors*, vol. 19, no. 12, 2019.
- [8] F. Kunz et al., “Autonomous driving at Ulm University: A modular, robust, and sensor-independent fusion approach,” in *Proc. IEEE Intell. Veh. Symp. (IV)*, 2015.
- [9] M. Buchholz et al., “Handling occlusions in automated driving using a multiaccess edge computing server-based environment model from infrastructure sensors,” *IEEE Trans. Intell. Transp. Syst.*, vol. 14, no. 3, pp. 106–120, 2022.
- [10] —, “Automation of the UNICARagil vehicles,” *Open Access Repository der Universität Ulm*, 2020, doi: 10.18725/OPARU-34024.
- [11] M. Herrmann, A. Piroli, J. Strohbeck, J. Müller, and M. Buchholz, “LMB filter based tracking allowing for multiple hypotheses in object reference point association,” in *Proc. IEEE Int. Conf. Multisens. Fusion Integr. (MFI)*, 2020, pp. 197–203.
- [12] B. N. Vo, B. T. Vo, and M. Beard, “Multi-Sensor Multi-Object Tracking with the Generalized Labeled Multi-Bernoulli Filter,” *IEEE Trans. Signal Process.*, vol. 67, no. 23, pp. 5952–5967, 2019.
- [13] K. Da, T. Li, Y. Zhu, H. Fan, and Q. Fu, “Recent advances in multisensor multitarget tracking using random finite set,” *Frontiers of Information Technology & Electronic Engineering*, vol. 22, no. 1, pp. 5–24, 2021-01.
- [14] B.-N. Vo, B.-T. Vo, and H. G. Hoang, “An efficient implementation of the generalized labeled multi-Bernoulli filter,” *IEEE Trans. Signal Process.*, vol. 65, no. 8, pp. 1975–1987, 2017.
- [15] B. Wei, B. Nener, W. Liu, and L. Ma, “Centralized multi-sensor multi-target tracking with labeled random finite sets,” in *Proc. Int. Conf. Control, Autom. Inf. Sci. (ICCAIS)*, 2016, pp. 82–87.

- [16] S. Robertson, C. van Daalen, and J. du Preez, "Efficient approximations of the multi-sensor labelled multi-Bernoulli filter," *Signal Processing*, vol. 199, 2021.
- [17] M. Herrmann, C. Hermann, and M. Buchholz, "Distributed implementation of the centralized generalized labeled multi-Bernoulli filter," *IEEE Trans. Signal Process.*, vol. 69, pp. 5159–5174, 2021.
- [18] M. Herrmann, T. Luchterhand, C. Hermann, T. Wodtke, J. Strohbeck, and M. Buchholz, "Notes on the product multi-sensor generalized labeled multi-Bernoulli filter and its implementation," in *Proc. Int. Conf. Inf. Fusion (FUSION)*, 2022.
- [19] M. Herrmann, T. Luchterhand, C. Hermann, and M. Buchholz, "The product multi-sensor labeled multi-Bernoulli filter," in *Proc. Int. Conf. Inf. Fusion (FUSION)*, 2023.
- [20] C. Hermann, M. Herrmann, T. Griebel, and M. Buchholz, "The fast product multi-sensor labeled multi-Bernoulli filter," in *Proc. Int. Conf. Inf. Fusion (FUSION)*, 2023.
- [21] S. Reuter, A. Danzer, M. Stubler, A. Scheel, and K. Granstrom, "A fast implementation of the Labeled Multi-Bernoulli filter using Gibbs sampling," in *2015 IEEE Intell. Veh. Symp. (IV)*, 2017, pp. 765–772.
- [22] C. Shim, B.-T. Vo, B.-N. Vo, J. Ong, and D. Moratuwage, "Linear complexity gibbs sampling for generalized labeled multi-bernoulli filtering," *IEEE Trans. Signal Process.*, vol. 71, pp. 1981–1994, 2023.
- [23] M. Beard, B. T. Vo, and B. Vo, "A solution for large-scale multi-object tracking," *IEEE Trans. Signal Process.*, vol. 68, pp. 2754–2769, 2020.
- [24] R. Kruse, E. Schwecke, and J. Heinsohn, *Uncertainty and Vagueness in Knowledge Based Systems*. Springer Berlin, Heidelberg, 1991.
- [25] OpenStreetMap contributors. (2017) Planet dump. Retrieved from <https://planet.osm.org>.
- [26] F. Poggenghans et al., "Lanelet2: A high-definition map framework for the future of automated driving," in *Proc. IEEE Intell. Transp. Syst. Conf.*, 2018.
- [27] J. Strohbeck, J. Muller, A. Holzbock, and M. Buchholz, "DeepSIL: A software-in-the-loop framework for evaluating motion planning schemes using multiple trajectory prediction networks," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst. (IROS)*, 2021.
- [28] M. Beard, B. T. Vo, and B.-N. Vo, "OSPA(2): Using the OSPA metric to evaluate multi-target tracking performance," in *Proc. Int. Conf. Control, Autom. Inf. Sci. (ICCAIS)*, 2017, pp. 86–91.
- [29] A. F. Garcia-Fernandez, L. Svensson, J. L. Williams, Y. Xia, and K. Granstrom, "Trajectory Poisson multi-Bernoulli filters," *IEEE Trans. Signal Process.*, vol. 68, pp. 4933–4945, 2020.

Physics-informed Reinforcement Learning for Automated Merging in Dense Traffic

Johannes Fischer*, Alexei Trofimov† and Christoph Stiller‡

Abstract: Decision-making in interactive traffic situations is a challenging task for automated vehicles. Reinforcement learning (RL) is a promising approach to learn a driving policy from interactions with a simulator or from real driving data. However, reinforcement learning often requires many interactions with the environment and can have difficulties with generalization to unseen situations. To resolve these problems, we propose to use physics-informed deep learning to regularize the RL algorithm with a driver model. In our evaluation we show that this approach leads to improved sample efficiency and better generalization to more challenging scenarios.

Keywords: Driver Model, Merging, Physics-informed, Deep Learning, Reinforcement Learning

1 Introduction

In highly complex and interactive traffic situations it is difficult for automated vehicles to make optimal decisions. Previous work has used Reinforcement Learning (RL) to learn a cooperative policy for navigating interactive traffic situations [1]–[4]. Other works have used imitation learning to learn human-like driving behavior [5], [6].

A well-known problem with reinforcement learning is that it requires many interactions with the environment to learn a policy. Moreover, it can exhibit bad generalization when used in domains that are different from the training domain. At last, it can also be difficult for the RL algorithm to converge to the optimal policy.

In this work, we introduce a new approach to improve the sample efficiency and generalization of RL algorithms. Based on the principles of Physics-Informed Deep Learning (PIDL), we regularize a policy gradient algorithm with a physics model that approximately follows the desired behavior. PIDL is known to improve sample efficiency and generalization [7]. Furthermore, it can also help guide the policy towards the desired behavior. [8] has used PIDL with behavior cloning for automated driving.

We apply the resulting algorithm to an interactive merging scenario. For this reason, we use the Gap Approaching Intelligent Driver Model (GAP-IDM) [9] as the physical model. The GAP-IDM is a driver model designed to smoothly approach traffic gaps while considering distances to multiple vehicles. In our evaluation, we consider different traffic

*Johannes Fischer, M.Sc., is PhD student at the Institute of Measurement and Control Systems (MRT) at Karlsruhe Institute of Technology (KIT), johannes.fischer@kit.edu.

†Alexei Trofimov, B.Sc., is student at Karlsruhe Institute of Technology (KIT)

‡Prof. Dr.-Ing. Christoph Stiller is head of the Institute of Measurement and Control Systems (MRT) at Karlsruhe Institute of Technology (KIT), stiller@kit.edu.

conditions on the target lane. In particular, the policy has to navigate environments with more challenging traffic densities than the training environment.

2 Technical background

In this section, we will give a brief introduction to the reinforcement learning problem in fully and partially observable environments, the physics-informed deep learning approach and the Gap Approaching Intelligent Driver Model.

2.1 Reinforcement Learning

Sequential decision-making problems can be described as **Markov Decision Processes (MDPs)**, where the environment's state s_t can be influenced by actions a_t at discrete time steps t . The decision-making agent follows a **policy**, which is a mapping from states to a probability distribution over actions. After each step, the agent receives a reward $\mathcal{R}(s_t, a_t)$ and the environment stochastically transitions to a new state s_{t+1} depending on the current state and the action. The goal is to find the policy π that maximizes the sum of cumulative discounted rewards $\sum_{t=0}^{\infty} \gamma^t \mathcal{R}(s_t, a_t)$ [10].

Reinforcement Learning (RL) tackles this problem by learning a policy π through interactions with the environment. **Proximal Policy Optimization (PPO)** is an online policy gradient method that strikes a balance between exploration and exploitation while effectively constraining policy updates become more sample-efficient [11].

2.2 Partially Observable Environments

In a **Partially Observable Markov Decision Process (POMDP)**, the agent does not have access to the environment's state s_t but only to noisy observations o_t . To make optimal decisions, the agent has to infer a probability distribution over the environment's state $p(s_t | a_{0:t-1}, o_{1:t})$ given the history of previous actions and observations. This distribution is called the belief b_t and is the input for a policy in a POMDP. Every POMDP can be interpreted as an MDP with beliefs b as states [12]. This so-called **belief MDP** can also be solved using RL algorithms [1].

2.3 Physics-informed Deep Learning

At the intersection between data-driven methods and model-based methods, **Physics-Informed Deep Learning (PIDL)** emerges as an approach that integrates physical knowledge into the training process of neural networks. To this end, the loss function is augmented with a term that enforces physical constraints [7].

In a standard supervised deep learning task a neural network ϕ_θ is regressed on labelled data (x_i, y_i) with the mean squared error loss $L(\theta) = \sum_i (\phi_\theta(x_i) - y_i)^2$.

If the data is known to satisfy a functional relation $f(x, y) = 0$, PIDL can be used to enforce this relation by augmenting the loss function with a term that penalizes the violation of this relation. The loss function is then given by $L(\theta) = \sum_i [(\phi_\theta(x_i) - y_i)^2 + \lambda f(x_i, y_i)^2]$, where λ is a hyperparameter to weigh the regularization with the physical relation.

2.4 Gap Approaching Intelligent Driver Model

Simple car-following driver models like the Intelligent Driver Model (IDM) [13] face challenges in interactive scenarios like lane changes or merging [9]. To address these issues, the GAP-IDM extends the IDM by considering multiple front and rear target vehicles and enabling smoothly approaching traffic gaps even when the vehicle is not initially aligned with the gap [9]. The acceleration of the GAP-IDM is given by

$$a_{\text{GAP-IDM}}(v, s_f, v_f, s_r, v_r) = a_{\max} \cdot \left(1 - \left(\frac{v}{v_{\text{des}}} \right)^4 - \left(\frac{s^*(v, v_f)}{g(s_f)} \right)^2 + \left(\frac{s^*(v_r, v)}{g(s_r)} \right)^2 \right) \quad (1)$$

with the desired distance

$$s^*(v, v_f) = s_{\text{des}} + \max \left(0, vT + \frac{v(v - v_f)}{2\sqrt{a_{\max} \cdot d_{\text{cmf}}} \right)$$

where v is the ego velocity, s_f, v_f, s_r, v_r are the signed distance and velocity of the most relevant front and rear target vehicles, respectively, g is a distance rectifier and the parameters $(a_{\max}, d_{\text{cmf}}, v_{\text{des}}, s_{\text{des}}, T)$ are the maximum acceleration, comfortable deceleration, desired velocity, minimum desired distance, and time headway. For this work, we use the shifted soft plus rectifier $g_{\alpha, \beta}(s) = \frac{1}{\beta} \log(1 + \alpha + \exp(\beta s))$ with a sharpness parameter $\beta > 0$ and a shifting parameter $\alpha \geq 0$.

3 Autonomous Merging Problem

In this section we describe the merging scenario that is used for evaluating our automated merging approach. It represents a highway on-ramp situation with dense traffic on the main roadway, where the merging vehicle has to find a suitable gap and assess the cooperation of other drivers [1]. While human drivers can assess and manage such a situation through experience and their individual driving style, it is extremely challenging for an automated system to learn intelligent behavior in this scenario. In this section, we provide an overview on how the merging scenario is modeled. A more detailed description can be found in prior work [1], [4].

3.1 Environment Description

The environment, as visualized in Fig. 1, consists of a merge lane where the agent is placed, and a main lane with dense traffic where each vehicle has a varying level of cooperativeness, which impacts their behavior. The behavior of the vehicles on the main lane is modeled by the cooperative IDM (C-IDM) [1]. This means, they will generally follow IDM behavior with respect to the vehicle in front. Additionally, they will yield to the merging vehicle based on the time to reach the merge point (TTM) if $\text{TTM}_{\text{merge}} < c \cdot \text{TTM}_{\text{main}}$, where $c \in [0, 1]$ is their cooperation parameter. That is, if the merging vehicle is expected to reach the merge point before the main lane vehicle reaches the merge point, weighted by a factor of c . As a consequence, the agent has to show its merging intent for other vehicles to react and yield.

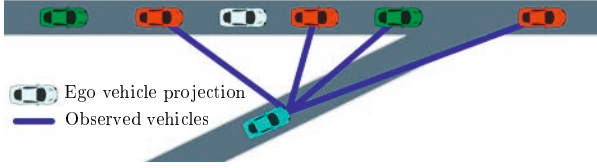


Figure 1: The ego vehicle (cyan) has to merge onto the main road where cooperative vehicles (green) might yield while non-cooperative vehicles (red) ignore it. The ego vehicle observes the vehicles most relevant for merging.

3.2 Agent Modeling

The agent observes its own physical state and the physical state of the four most relevant vehicles within its field of view. These vehicles are the vehicles before and after the merge point and before and after the projection of the ego vehicle onto the main lane, as illustrated in Fig. 1. The cooperation levels of the vehicles are not observable. Therefore, the agent must infer a belief over the cooperation to perform the merging maneuver. In this work, we use the Bayesian filter for inferring the belief that was introduced in prior work [1]. This modeling approximates the cooperation belief as a Bernoulli distribution for each vehicle, resulting in a low-dimensional belief state.

At each discrete time step of $\Delta t = 1$ s, the agent can choose between three jerk levels $\{1 \text{ m/s}^3, 0 \text{ m/s}^3, -1 \text{ m/s}^3\}$. At each time step the agent is rewarded a positive reward of +100 if reaching the goal behind the merge point, a negative reward of -100 if colliding with another vehicle, and a negative reward of $-0.1 \cdot (a_t^2 + j_t^2)$ for making use of acceleration a_t or jerk j_t . An episode terminates after reaching the goal, after a collision or after 100 time steps.

4 Physics-informed Reinforcement Learning

In our approach, we use a driver model as a physical equation to regularize a reinforcement learning policy. We begin with describing how the GAP-IDM is used as a physical equation in the merging scenario and then illustrate how this equation is used in the reinforcement learning algorithm.

4.1 GAP-IDM as Regularization

To use the GAP-IDM in a PIDL architecture, it needs to be reformulated as a function that maps the agent's observation to an acceleration. GAP-IDM is designed to output accelerations to approach a target gap, but does not decide on which gap to target. Hence, the target gap needs to be determined from the observation beforehand.

To this end, we use a neural network to predict the target gap from the observation. We formulate this problem as a classification task, where the model decides between four possible target gaps. The first gap is always the one behind the vehicle directly in front of the merging point. The other three gaps are the subsequent gaps on the main lane, as illustrated in Fig. 2.

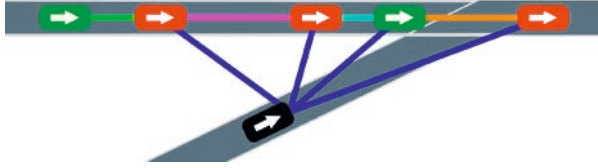


Figure 2: The four gaps considered by the ego vehicle (orange, cyan, magenta and green).

To produce human-like behavior, the classifier should be fit to human driving data. In this work, we use driving data generated with a trained PPO policy as a surrogate for human driving data. To ensure that this results in a good policy, we explicitly provide the cooperation levels of observed vehicles as an additional input to the agent. The generated data is used to train the gap classifier based on the gap chosen by the PPO policy.

4.2 Physics-informed Proximal Policy Optimization

The idea is to use the trained gap classifier to select the target vehicles for the GAP-IDM and use the GAP-IDM to regularize the reinforcement learning algorithm to produce actions that result in a behavior closer to the GAP-IDM. In this work, we will use the PPO algorithm, but the idea can be used with other policy gradient algorithms in the same way.

We incorporate the physics model $f(s, a) = 0$ into the PPO loss by an additional loss term with a regularization weight λ . For a mini-batch of states $\mathcal{B}$ and the problem's action space $\mathcal{A}$, the loss term is defined as

$$L^{\text{phy}} = \frac{1}{|\mathcal{B}|} \sum_{s \in \mathcal{B}} \left(\frac{1}{|\mathcal{A}|} \sum_{a \in \mathcal{A}} \pi(a|s) f(s, a)^2 \right). \quad (2)$$

This residual loss term is minimized by increasing the probability $\pi(a|s)$ of selecting actions with a low mean squared model error and decreasing the probability of selecting actions with a high mean squared model error. The weight λ can be decreased during training to initially use the physics model as guidance but gradually recover the original PPO objective. We refer to the resulting algorithm as Physics-Informed Proximal Policy Optimization (PI-PPO).

4.3 PI-PPO in the Merging Scenario

To apply PI-PPO to the merging scenario, we use the error between the acceleration resulting from the chosen jerk action and the predicted acceleration of the GAP-IDM as the physics model. Furthermore, the environment is only partially observable. For this reason, we use the belief b as input to the policy, which contains the physical state observations and the inferred cooperation beliefs. Therefore, the GAP-IDM loss for the partially observable merging environment is given as

$$L^{\text{phy}} = \frac{1}{|\mathcal{B}|} \sum_{b \in \mathcal{B}} \left(\frac{1}{|\mathcal{A}|} \sum_{j \in \mathcal{A}} \pi(j|b) (\text{acc}^e + j \cdot \Delta t - \text{acc}^{\text{phy}}(b))^2 \right) \quad (3)$$

Parameter	Traffic condition	
	Moderate	Dense
$N_{\min}$	4	8
$N_{\max}$	8	12
p_{spawn}	1.0	0.3
$v_{\text{des},\min}$	4 m/s	4 m/s
$v_{\text{des},\max}$	6 m/s	6 m/s

Table 1: Scenario-specific parameters.

where $\mathcal{A}$ is the action space of available jerk levels j , $\mathcal{B}$ is a mini-batch of beliefs b , acc^e is the current ego acceleration, Δt is the time step, and $\text{acc}^{\text{phy}}(b)$ is the predicted acceleration of the GAP-IDM.

5 Experiments and Evaluation

We evaluate our approach on the merging scenario described in Section 3. To assess how well our algorithm generalizes to unseen situations, we consider different traffic conditions on the target lane. The traffic scenarios vary in the uniformly sampled number of vehicles on the main lane N , the probability that vehicles re-enter the main-lane after reaching its end p_{spawn} , and their uniformly sampled desired speed v_{des} according to Table 1. The vehicles on the main lane are simulated according to C-IDM with parameters $a_{\max} = 2 \text{ m/s}^2$, $d_{\text{cmf}} = 2 \text{ m/s}^2$, $s_{\text{des}} = 2 \text{ m}$, $T = 1.5 \text{ s}$. All agents are trained in the moderate traffic scenario and evaluated on more dense traffic compared to the training environment.

To train the gap classifier, an oracle agent with full knowledge of the latent cooperation levels of other vehicles is trained in the moderate traffic scenario. This is done to ensure that the data used to train the gap classifier is of high quality. The gap classifier is also only trained on data collected in moderate traffic.

The training parameters for PPO and PI-PPO were separately optimized using random search and are provided in Table 2. The parameters for the GAP-IDM used in PI-PPO are the same as the C-IDM parameters, except for the desired speed: We choose $v_{\text{des}} = 15 \text{ m/s}$ larger than the desired speed of the vehicles on the main lane to ensure that the agent is not incentivized to slow down before merging. The parameter values of the shifted softplus rectifier g are chosen as $\alpha = 5$ and $\beta = 0.3$.

The training progress in the moderate traffic scenario is depicted in Fig. 3. The physics-informed algorithm can be seen to converge faster compared to the uninformed PPO algorithm since the GAP-IDM loss term is able to guide the policy towards reasonable behavior. This illustrates the improved sample efficiency of the physics-informed algorithm.

To minimize deviation due to the random initialization of the traffic scene, each algorithm is evaluated on 1000 episodes in the test settings. Table 3 shows the results of the evaluation in moderate and dense traffic, measured by the average episode return, the success rate and the average episode length. Unsuccessful episodes are those that result in a collision. The performance of PI-PPO and PPO is similar on the training environ-

Parameter	Value
Neural network architecture	3 dense layers, (128, 128, 64) nodes
Activation function	ELU (except for output layer)
Epochs per batch	8
Optimizer	Adam [14]
Learning rate	$8 \cdot 10^{-4}$
Batch size	800
Training steps	$1 \cdot 10^6$
Discount factor γ	0.95
Clip ratio	0.15
Critic loss weight	0.5
Entropy regularization weight	$1 \cdot 10^{-3}$ (PI-PPO), $8 \cdot 10^{-3}$ (PPO)
Physics loss weight λ	0.2
Scheduling for λ	Linear to zero in the first 30% training steps

Table 2: Parameters used for training the PPO and PI-PPO agents.

Traffic condition	Algorithm	Episode return	Success rate [%]	Episode length
Moderate	PI-PPO	88.7 	95.8 	16.0 
	PPO	87.2 	95.3 	16.1 
Dense	PI-PPO	85.4 	94.5 	18.5 
	PPO	69.0 	86.4 	17.6 

Table 3: Evaluation results on moderate and dense traffic scenarios.

ment with moderate traffic density. However, on the more challenging test environments, PI-PPO outperforms PPO in terms of average return and success rate.

6 Conclusions and Future Work

In this work, we present a framework for physics-informed reinforcement learning in an automated merging scenario. We use the GAP-IDM as a driver model to regularize the policy gradient algorithm PPO with a policy loss term that penalizes deviations from the driver model. As our evaluation shows this leads to improved convergence properties, higher sample efficiency, and better generalization abilities to unseen traffic conditions compared to the physics-uninformed algorithm.

In our future research we want to investigate to which extent physics-informed deep learning can be used in the context of imitation learning. Since algorithms like Adversarial Inverse Reinforcement Learning employ a policy as the generator [15], PI-PPO could be used as a drop-in replacement. This could lead to faster and more stable training. Another interesting possibility is to test the algorithm on real driving data. This includes fitting the gap classifier and the driver model to real driving data and then training PI-PPO on this data.

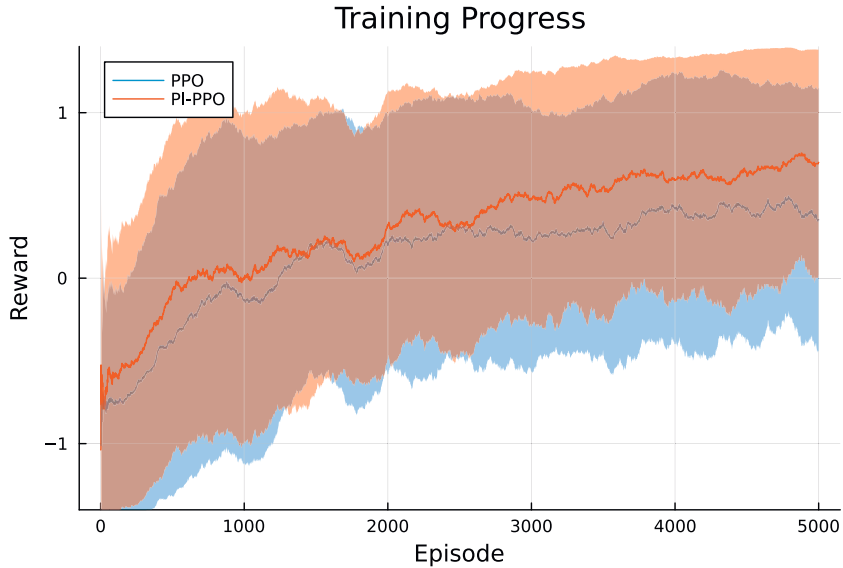


Figure 3: Moving average of the reward per episode during training with a sliding window size of $w = 250$ episodes for the PPO and PI-PPO algorithms. The darker line represents the moving average and the lighter line the moving standard deviation. Only the first 5000 episodes of training are shown.

References

- [1] M. Bouton, A. Nakhaei, K. Fujimura, and M. J. Kochenderfer, “Cooperation-Aware Reinforcement Learning for Merging in Dense Traffic,” in *2019 IEEE Intelligent Transportation Systems Conference (ITSC)*, Oct. 2019, pp. 3441–3447.
- [2] D. Kamran, T. Engelgeh, M. Busch, J. Fischer, and C. Stiller, “Minimizing Safety Interference for Safe and Comfortable Automated Driving with Distributional Reinforcement Learning,” in *2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, Prague, Czech Republic, Sep. 2021, pp. 1236–1243.
- [3] M. Sackmann, H. Bey, U. Hofmann, and J. Thielecke, “Learning a Diverse and Cooperative Policy for Predicting Roundabout Traffic Situations,” in *14. Uni-DAS e.V. Workshop Fahrerassistenz Und Automatisiertes Fahren*, 2022-05-09/2022-05-11, 2022.
- [4] J. Fischer, E. Bührle, D. Kamran, and C. Stiller, “Guiding Belief Space Planning with Learned Models for Interactive Merging,” in *2022 IEEE 25th International Conference on Intelligent Transportation Systems (ITSC)*, Macau, China, Oct. 2022, pp. 2542–2549.

- [5] J. Fischer, C. Eyberg, M. Werling, and M. Lauer, “Sampling-based Inverse Reinforcement Learning Algorithms with Safety Constraints,” in *2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, Prague, Czech Republic, Sep. 2021, pp. 791–798.
- [6] M. Sackmann, H. Bey, U. Hofmann, and J. Thielecke, “Modeling Driver Behavior using Adversarial Inverse Reinforcement Learning,” in *2022 IEEE Intelligent Vehicles Symposium (IV)*, Jun. 2022, pp. 1683–1690.
- [7] M. Raissi, P. Perdikaris, and G. Karniadakis, “Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving non-linear partial differential equations,” *Journal of Computational Physics*, vol. 378, pp. 686–707, Feb. 2019, issn: 00219991.
- [8] Z. Mo, X. Di, and R. Shi, “A Physics-Informed Deep Learning Paradigm for Car-Following Models,” *Transportation Research Part C: Emerging Technologies*, vol. 130, p. 103240, Sep. 2021, issn: 0968090X. arXiv: 2012.13376 [cs, eess].
- [9] J. Fischer, E. Bührle, and C. Stiller, “Gap Approaching Intelligent Driver Model for Interactive Simulation of Merging Scenarios,” in *2023 IEEE Intelligent Vehicles Symposium (IV)*, Anchorage, United States, Jun. 2023, pp. 1–8.
- [10] R. S. Sutton and A. Barto, *Reinforcement Learning: An Introduction*, Second edition, ser. Adaptive Computation and Machine Learning. Cambridge, MA London: The MIT Press, 2018, ISBN: 978-0-262-03924-6.
- [11] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, “Proximal Policy Optimization Algorithms,” *arXiv:1707.06347 [cs]*, Aug. 2017. arXiv: 1707.06347 [cs].
- [12] M. Kochenderfer, *Decision Making Under Uncertainty - Theory and Application*, 1st. MIT Press, 2015, ISBN: 0-262-02925-1 978-0-262-02925-4.
- [13] M. Treiber, A. Hennecke, and D. Helbing, “Congested Traffic States in Empirical Observations and Microscopic Simulations,” *Physical Review E*, vol. 62, no. 2, pp. 1805–1824, Aug. 2000, issn: 1063-651X, 1095-3787.
- [14] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” in *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, Y. Bengio and Y. LeCun, Eds., 2015.
- [15] J. Fu, K. Luo, and S. Levine, “Learning robust rewards with adversarial inverse reinforcement learning,” in *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*, 2018.

Reducing Ghost Detections Through Uncertainty Modeling for Automated Driving

Ahmed Hammam^{*‡} Frank Bonarens^{*},
Seyed Eghbal Ghobadi[†] und Christoph Stiller[‡]

Abstract: Deep neural networks (DNN) have demonstrated remarkable performance in various tasks related to automated driving. However, one significant obstacle hindering their application in automated driving systems is the occurrence of false positive detections. In our context, false positive detections are referred to as ghost detections, wherein the DNN mistakenly identifies parts of the scene as objects. In this work, we explore the prospect of leveraging uncertainty modeling to effectively minimize ghost detections. We propose a method that builds on an instance segmentation framework that better separates true positive from false positive distributions than state-of-the-art methods. This method integrates the Intermediate Layer Variational Inference (ILVI) approach and Dirichlet Distributions into an instance segmentation network. Our experimental results demonstrate that our proposed method not only enhances instance and semantic segmentation performance but also improves uncertainty estimation. Leveraging significantly improved uncertainty estimation, we investigate the potential of thresholding on uncertainty to reduce the occurrence of ghost detections, thereby enhancing both precision and recall performance.

Keywords: Deep Neural Networks, Uncertainty Estimation, Instance Segmentation

1 Introduction

Deep learning has revolutionized computer vision, offering groundbreaking advancements in various domains including medical imaging [1] and automated driving (AD) [2]. In the field of AD systems, deep neural networks (DNNs) have emerged as the predominant method, finding widespread applications in sensor fusion [3], path planning [4] and image semantic segmentation [5].

Even though DNNs deliver high performance on their trained tasks, DNNs are overconfident by delivering unreliable high confidence on incorrect predictions [6]. This drawback becomes a serious limitation for AD systems. This typical insufficiency of a DNN leads to false or ghost detections, reducing the overall performance of an AD system [7].

In recent years, a common solution to this limitation involves enhancing DNNs with the capability to explicitly express uncertainty regarding their output predictions [8, 9]. The introduction of uncertainty modeling stands as a fundamental advancement in mitigating

^{*}Stellantis, Opel Automobile GmbH, Rüsselsheim am Main.
(ahmedmostafa.hammam@external.stellantis.com)

[†]Technische Hochschule Mittelhessen, Gießen.

[‡]Institute of Measurement and Control Systems, Karlsruhe Institute of Technology, Karlsruhe.

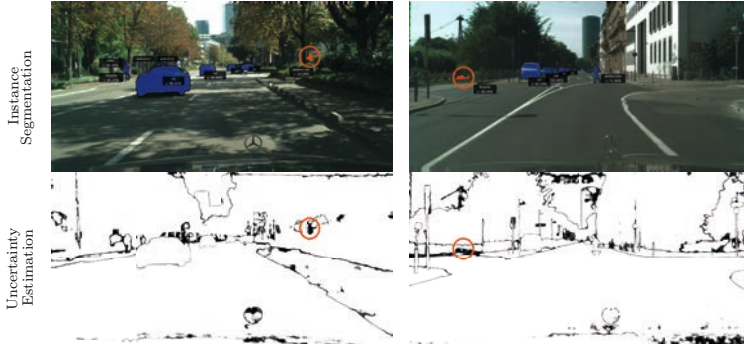


Figure 1: Sample DNN Outputs: Instance segmentation results (Top Row) and corresponding uncertainty estimates (Bottom Row). True positive detections are color-coded in blue, denoting high certainty; false detections are color-coded in red, indicating low certainty. The DNN effectively distinguishes between both detections, associating high certainty to true positives and low certainty to false ghost detections.

the insufficiencies of DNNs, particularly their tendencies toward overconfidence or lack of precision [10].

Building upon prior research [11–13], our study introduces a method centered around an enhanced architecture incorporating two key layers: Dirichlet Maximum Likelihood Estimation (MLE) and Intermediate Layer Variational Inference (ILVI). The integration of these two layers is to empower the DNN for accurate instance segmentation within scenes, while concurrently ensuring reliable uncertainty estimation.

By achieving reliable uncertainty estimates, our approach addresses the issue of ghost detections by being able to distinguish between true positives from false positives. This separation enables the incorporation of thresholding techniques, which play a crucial role in reducing ghost detections. Consequently, our methodology enhances uncertainty estimation capabilities and contributes to more accurate and robust results.

The remainder of the paper is organized as follows: Section 2 presents related work to Dirichlet modeling and its application in uncertainty estimation, whilst the architecture is discussed and explained in Section 3. The experiments conducted to test our approach are displayed in Section 4 and a conclusion of the work is discussed in Section 5.

2 Related Work

In recent years, efforts have concentrated on enhancing the credibility of generated outputs by modeling uncertainty within DNNs [14, 15]. Specifically, the focus has shifted towards modeling DNN outputs as Dirichlet distributions to refine uncertainty estimation. One approach has been Dirichlet Prior Networks, which builds upon the framework introduced in [16] by modeling the predicted logits from the DNN as the concentration parameters of a Dirichlet distribution that serves as a prior for the categorical distribution. Similar to Prior networks, the authors in [17] proposed a method that combines Prior networks with

the likelihood to maximize the whole posterior. Inspired by the Dempster-Shafer theory of evidence (DST) [18], they treat the predictions of the DNN as subjective opinions and train the DNN to gather evidence supporting these opinions. Additionally, a penalty term is introduced to penalize the DNN for incorrect detections and encourage it to exhibit high uncertainty in such cases.

Inspired by previous studies of [16, 17], our goal is to utilize Dirichlet models to enhance the reliability of uncertainty estimation and maintain segmentation performance. Optimizing the reliability of uncertainty estimation in the Dirichlet DNN by formulating its loss function using KL divergence is often considered challenging [16].

3 Methodology

In this section, we outline the core components of our architecture, presented in Figure 2, emphasizing their distinct roles. Subsequently, we explore the semantic segmentation decoder in-depth, showcasing its integration with the Dirichlet layer for improved uncertainty estimation. We then explain the ILVI approach, followed by a description of the applied thresholding methodology aimed at mitigating false positive occurrences.

3.1 Dirichlet DNN Architecture

The architecture, presented in Figure 2a, comprises a shared backbone that takes the input image and passes the extracted features to the ILVI module. The ILVI module, presented in Figure 2b, acts as a regularizer by adding stochasticity in the DNN avoiding overfitting and overconfidence. The output of the ILVI is passed on to the semantic segmentation decoder and the instance segmentation decoder.

The semantic segmentation decoder and the Dirichlet layer, shown in Figure 2c, are trained together to model the semantic segmentation output as a Dirichlet distribution, which enables the uncertainty estimation. This decoder generates two results; semantic segmentation and uncertainty estimation based on the per-pixel Dirichlet distribution. The instance segmentation decoder generates the center points and the masks of the instances. The generated outputs from both decoders are passed onto the post-processing module. For every instance identified by the instance decoder, the instance mask, instance class, and the average uncertainty estimate are provided.

The architecture integrates the lightweight MobileNetV3 [19] as its foundational backbone and encompasses semantic and instance segmentation decoders influenced from Panoptic Deeplab [20].

The training of this architecture is governed by the following composite loss function:

$$\mathcal{L} = \mathcal{L}_{sem} + \mathcal{L}_{ILVI} + \mathcal{L}_{ins} \quad (1)$$

Here, $\mathcal{L}_{sem}$ represents the loss function for per-pixel classification in semantic segmentation, combined with the Dirichlet distribution modeling. $\mathcal{L}_{ILVI}$ stands for the ILVI loss which introduces stochasticity to the DNN, thereby enhancing its uncertainty estimation capacity and generalization efficacy. The third component, $\mathcal{L}_{ins}$, pertains to the instance segmentation loss. The specifics of each loss term are further explained in the corresponding following sections. [21]

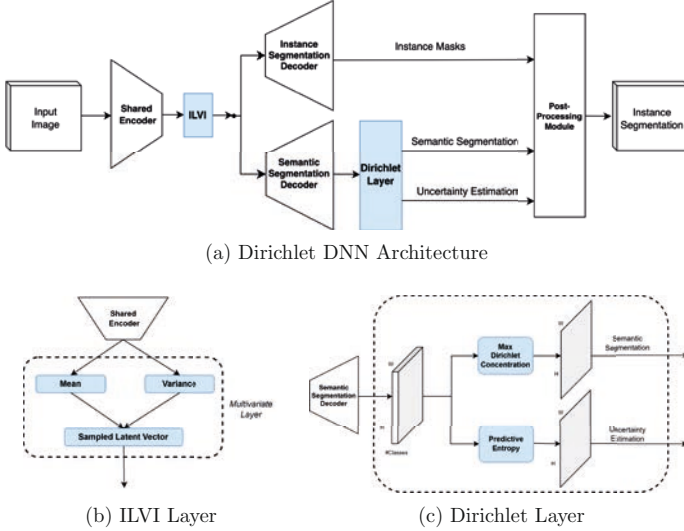


Figure 2: The Dirichlet MLE DNN is illustrated, emphasizing its key components for enhancing OOD identification. The ILVI layer introduces a multivariate layer structure, while the Dirichlet layer handles semantic segmentation and uncertainty estimation.

3.2 Semantic Segmentation Decoder

In this section, we explain the fundamental concepts and methods at the core of our approach. We begin with an in-depth exploration of the Dirichlet distribution within the probability simplex, a foundational probabilistic structure characterized by concentration parameters.

A supervised network aims to predict the target value $y \in \mathcal{Y}$ for an input $x \in \mathcal{X}$, where the input space $\mathcal{X}$ corresponds to the space of images. Accordingly, a supervised machine learning problem with the task of semantic segmentation has a target $\mathcal{Y}$ consisting of a finite set of c classes where the task for the network is to predict the class of each pixel out of the set of classes K . For our purpose, a DNN is defined as a function $f_w : \mathcal{X} \rightarrow \mathcal{Y}$, parameterized by $w \in \mathbb{R}$, which maps an input $x \in \mathcal{X}$ to an output $f_w(x) \in \mathcal{Y}$.

Given the probability simplex as $\mathcal{S} = \{(\theta_1, \dots, \theta_k) : \theta_i \geq 0, \sum_i \theta_i = 1\}$, the Dirichlet distribution is a probability density function on vectors $\theta \in \mathcal{S}$ and categorized by concentration parameters $\alpha = \{\alpha_1, \dots, \alpha_K\}$ as:

$$\text{Dir}(\theta; \alpha) = \frac{1}{B(\alpha)} \prod_{i=1}^K \theta_i^{\alpha_i-1} \quad (2)$$

where the normalizing constant $\frac{1}{B(\alpha)}$ denotes the multivariate Beta function $B(\alpha) = \frac{\prod_{i=1}^K \Gamma(\alpha_i)}{\Gamma(\alpha_0)}$, $\alpha_0 = \sum_{i=1}^K \alpha_i$ and Gamma function $\Gamma(x) = \int_0^\infty t^{x-1} e^{-t} dt$, and θ denotes the ground truth probability distribution [22]. To model the Dirichlet distribution, the con-

centration parameters α correspond to each class output from the semantic segmentation decoder as follows: $\alpha = f_w(x)$, where α changes with each input x .

To train the Dirichlet distributions, we propose a direct maximization of the likelihood, inspired by the works of [12, 13]. Unlike the Dirichlet Prior and Evidential approaches, our method eliminates the constraints of the KL-divergence term in the loss function, allowing the DNN to explore the weight space more freely. This leads to improved segmentation performance and enhanced reliability in uncertainty estimation by encouraging a sharper concentration of Dirichlet parameters for correct predictions and flatter distributions for incorrect predictions.

Training a Dirichlet DNN with maximum likelihood estimation (MLE) can be done by minimizing the negative log-likelihood [22] as follows:

$$F(\alpha; \theta) = \log \prod \text{Dir}(\theta; \alpha) = \log \Gamma \left(\sum_{i=1}^K \alpha_i \right) - \sum_{i=1}^K \log \Gamma(\alpha_i) + \sum_{i=1}^K (\alpha_i - 1) \log \theta_i \quad (3)$$

where θ represents the probability distribution to be maximized.

We aim to train the DNN to produce reliable uncertainty estimations by treating the DNN's correct and incorrect predictions separately. Our primary objective is to obtain accurate predictions with low uncertainties, while assigning high uncertainty to incorrect predictions.

To ensure high certainty for correct predictions, the DNN should exhibit a strong concentration toward the correct class, as shown in Figure 3a. This can be achieved by maximizing the likelihood using the ground truth label probability and employing a *one-hot vector*. Conversely, for incorrect predictions, high uncertainty is achieved by maximizing the likelihood using an equal probability vector with equal probabilities assigned to all classes, as shown in Figure 3b.

To address these cases, we extend the formulation presented in Equation 3 for the semantic segmentation as follows:

$$\mathcal{L}_{sem} = F(\alpha_{correct}; \theta_{correct}) + F(\alpha_{incorrect}; \theta_{incorrect}), \quad (4)$$

where $\alpha_{correct}$ and $\alpha_{incorrect}$ are the network's concentration parameters representing the correct and incorrect DNN predictions respectively, and $\theta_{correct}$ and $\theta_{incorrect}$ represent the ground truth probability distribution for the correct classes and the equal probability vector to yield high uncertainty respectively.

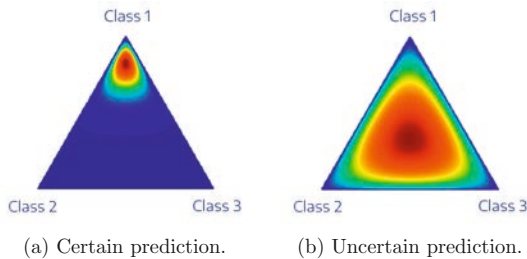


Figure 3: Dirichlet Plots

In this work, we model the uncertainty estimation of the Dirichlet distribution using the predictive entropy:

$$\hat{\mathbb{H}}[y|x] = - \sum_c (p(y = c|x, w)) \log(p(y = c|x, w)) \quad (5)$$

where y is the output variable, c ranges over all the classes K , $p(y = c|x, w) = \frac{\alpha_c}{\sum \alpha}$ is the probability of the input x being class c , and w are the model parameters. The class of each pixel for the semantic segmentation output is determined according to the highest concentration value of the Dirichlet distribution.

3.3 Intermediate Layer Variational Inference

The Intermediate Layer Variational Inference (ILVI) approach is designed to address the limitations of existing Bayesian deep neural network approximation techniques. The concept of ILVI, presented in Figure 2b, modifies a latent layer in the network to take the shape of a multivariate Gaussian distribution with mean and variance, instead of using single point estimates. Studies showed that by adopting this method, stochasticity is introduced allowing for the sampling of points from this layer and consequently improving the uncertainty estimation of the DNN and also its generalization performance [11].

The variational posterior of ILVI is modeled as a diagonal Gaussian distribution. The weight parameters w are sampled from this distribution using the reparametrization trick [23]. Specifically, each weight parameter is computed as: $w = \mu + \sigma \odot \epsilon$. Here, μ represents the mean of the Gaussian distribution, σ represents the standard deviation parameter, $\epsilon \sim \mathcal{N}(0, I)$ is a random variable, and $\odot$ is the pointwise multiplication. The parameter σ is pointwise parameterized as $\sigma = \log(1 + \exp(\rho))$, ensuring its non-negativity. This formulation preserves the mean and log-variance vectors as learnable parameters while introducing stochasticity through the random variable ϵ .

The ILVI method utilizes Bayesian variational inference based on the Kullback-Leibler (KL) divergence, following:

$$\mathcal{L}_{ILVI} = KL(q(\phi) || p(w|X, Y)). \quad (6)$$

This loss comprises of approximating a new probability distribution that is close to the posterior distribution produced by the model. To achieve the new approximate distribution, the KL divergence is needed to minimize the new variational parameters $q(\phi) \approx p(w|X, Y)$ for approximation.

3.4 Instance Segmentation

The instance segmentation decoder is a crucial component of the architecture, designed to identify individual object instances within an image. This decoder operates in conjunction with the semantic segmentation decoder to provide comprehensive scene understanding.

The output from the ILVI layer passes to the instance segmentation decoder to refine the features and generate predictions for object instances. Specifically, the decoder produces two outputs: instance masks and center points. Instance masks define the spatial boundaries of separate objects, while center points identify the most prominent pixel

within each object. These central points function as reference markers for precisely determining the locations of objects within the scene. Accordingly, the instance segmentation loss is formulated as follows:

$$\mathcal{L}_{ins} = \mathcal{L}_{center} + \mathcal{L}_{mask}. \quad (7)$$

The center points loss $\mathcal{L}_{center}$ employs the mean squared error loss to penalize the DNN’s center points with respect to the ground truth. In the case of the instance offset loss $\mathcal{L}_{mask}$, the L1 loss is utilized to penalize the difference between the DNN’s generated masks and the ground truth mask [20].

3.5 Thresholding for Uncertainty-Based Filtering

At its core, thresholding involves establishing a threshold value for the model’s uncertainty estimates. Detections with uncertainties below this threshold are discarded, effectively reducing false positive instances. The primary aim here is to enhance precision, the ratio of correctly predicted positive instances to all predicted positive instances. While this approach significantly improves precision, it may have effects on the recall, the ratio of correctly predicted positive instances to all actual positive instances. An inherent challenge with thresholding is the potential reduction in recall due to the elimination of detections below the threshold. This trade-off can result in missed true positive instances, which in turn may impact the model’s overall recall performance.

The effectiveness of thresholding relies on selecting the appropriate threshold value. To choose the threshold value, the precision and recall values should be plotted against varying threshold values. This visually illustrates the trade-off between these two crucial performance indicators. The threshold is selected to maximize precision improvement without compromising recall performance.

4 Experiments and Results

In the following section, we present a comprehensive analysis of our methodology and the corresponding results. Experiments conducted encompass evaluation of segmentation performance, uncertainty estimation, and uncertainty thresholding.

We undertake a thorough evaluation of our methodology’s performance in comparison to two state-of-the-art approaches: Prior Network, and Evidential Network whilst having Cross Entropy (CE) as our baseline. In this study, DNNs are trained on the Cityscapes dataset [24] and evaluated using its validation set. Additionally, we test their adaptability on the KITTI dataset [25], which examines the models’ resilience and real-world applicability across varying environments evaluating its generalization capabilities. In this work, we model the uncertainty estimation of all approaches using predictive entropy.

4.1 Segmentation Performance

The results of the DNNs’ segmentation performance are shown in Table 1. In this table, we present a comprehensive comparison of instance and semantic segmentation performance across the different approaches on both the Cityscapes and KITTI datasets. The metrics used include precision and precision up to 50 meters (Precision 50m) for

Table 1: Performance comparison for instance and semantic segmentation. The Dirichlet MLE + ILVI method achieves the highest scores across multiple metrics on both datasets, highlighting significant performance improvement.

	Cityscapes			KITTI	
	Precision	Precision 50m	mIoU	Precision	mIoU
CE	38.2	35.8	65.2	39.1	47.6
Prior	39.8	37.5	66.7	38.7	48.2
Evidential	42.9	39.8	68.1	40.6	50.1
Dirichlet MLE + ILVI	45.1	43.5	69.1	42.1	51.3

instance segmentation and mean Intersection over Union (mIoU) for semantic segmentation.

The effectiveness of the Dirichlet MLE + ILVI approach in our work is clearly demonstrated in the presented results in Table 1. Across both the Cityscapes and KITTI datasets, the Dirichlet MLE + ILVI approach consistently outperforms other techniques, showcasing its effectiveness in enhancing instance and semantic segmentation performance.

Moreover, the Dirichlet MLE + ILVI approach shows improvements in precision for the Cityscapes dataset, surpassing other methods. It also excels in precision 50%, precision 100m, and mIoU. This superiority extends to the KITTI dataset, with the highest precision 50% and mIoU. Consistent cross-dataset performance underscores its strength in instance and semantic segmentation. This highlights its potential to enhance localization accuracy and semantic understanding, making it invaluable for scene analysis tasks.

4.2 Uncertainty Estimation Performance

Two key evaluation metrics, namely separation efficiency and accuracy vs. certainty, are utilized to assess the efficacy of the proposed method in enhancing uncertainty estimation and prediction accuracy. Table 2 presents the results and offers a comprehensive insight into the performance of different uncertainty estimation approaches.

Distributional Separation Efficiency

We aim to quantify the efficiency of the DNN to differentiate between correct and incorrect predictions by plotting their corresponding certainty distribution for both cases. The distributions are then compared using the Wasserstein distance metric, where a high value indicates dissimilar distinctive distributions and vice versa. Notably, our Dirichlet MLE + ILVI method exhibits a significantly higher Wasserstein distance surpassing the other approaches.

Accuracy vs. Certainty

An important factor for deep neural networks is not only to be accurate about their predictions but also to be certain about them. Proposed by [26], ratios between accuracy and certainty are defined and quantified to compare between DNNs’ uncertainty performances.

Table 2: Separation efficiency and accuracy vs. certainty ($\uparrow$). These results collectively emphasize the efficacy of the Dirichlet MLE + ILVI method in achieving a balance between separation efficiency and accuracy across varying levels of certainty.

	Separation Efficiency	Accuracy vs. Certainty (%)		
	Wasserstein Distance	P(A C)	P(U I)	AvU
CE	1.1	50.9	24.4	51.8
Prior	2.3	72.7	17.5	38.5
Evidential	3.3	53.2	63.1	61.2
Dirichlet MLE + ILVI	4.9	85.4	78.1	70.1

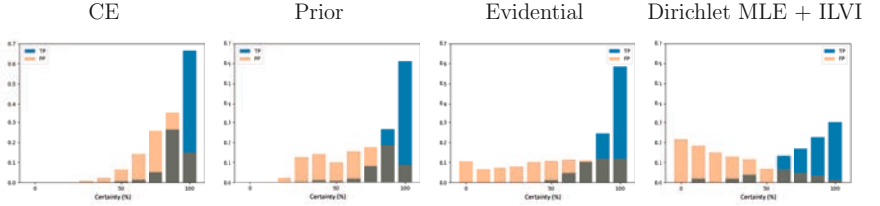


Figure 4: Distribution plots of true positives (TP) and false positives (FP) showing improved separation capability of Dirichlet MLE + ILVI method.

After attaining the DNN predictions and their respective uncertainty estimations, they are compared with the ground truth.

Three conditional probabilities are needed for this evaluation test, following the work of [26]: $p(\text{accurate}|\text{certain}) = \frac{n_{ac}}{n_{ac}+n_{ic}}$, $p(\text{uncertain}|\text{inaccurate}) = \frac{n_{iu}}{n_{ic}+n_{iu}}$ and $AvU = \frac{n_{ac}+n_{iu}}{n_{ac}+n_{au}+n_{ic}+n_{iu}}$, where n_{ac} are the accurate and certain predictions, n_{au} are the accurate and uncertain predictions, n_{ic} are the inaccurate and certain predictions, and n_{iu} are the inaccurate and uncertain predictions. AvU stands for accuracy vs. uncertainty, describing the probability of getting a good prediction out of the network either accurate and certain or inaccurate and uncertain.

In Table 2, our proposed method stands out with high accuracy vs. certainty percentage indicating the ability in generating predictions that are both accurate and certain. Additionally, it can be observed that there is a high improvement in the P(U|I) whilst still maintaining high performance on the other two metrics. This is not frequently observed as any method trying to improve uncertainty representation would come to a cost of reduced performance on the other two metrics. This finding underscores the approach’s capability to effectively recognize uncertain predictions.

4.3 Uncertainty Estimation Thresholding

By setting a certainty threshold, the DNN can effectively reduce the occurrence of false positives, leading to an improvement in its overall precision. The threshold acts as a cutoff value where detections less than the thresholded certainty are omitted.

For that, to choose the threshold we first plot the average precision and recall at varying thresholds, as shown in Figure 5. Average precision and recall are plotted for

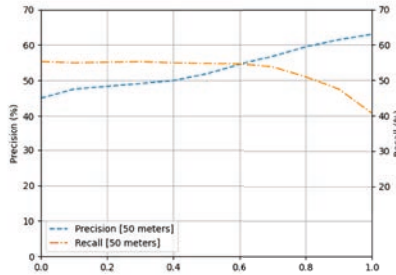


Figure 5: Average precision and recall values up to 50 meters are plotted with varying uncertainty thresholds.

Table 3: Precision and recall performance comparison before and after thresholding (in % ($\uparrow$)).

	Precision	Precision 50m	Recall	Recall 50m
CE	38.2	35.8	31.8	45.6
Prior	39.8	37.5	34.5	48.9
Evidential	42.9	39.8	37.1	51.3
Dirichlet MLE + ILVI	45.1	43.5	39.4	55.3
Dirichlet MLE + ILVI (Thresholded)	57.2	54.5	37.5	54.2

detections upto only 50 meters, as most false positive detections are observed within 50 meters range of the vehicle. This plot shows the precision and recall values at varying thresholds in steps of 10%. A threshold of 0% signifies the absence of thresholding, implying the utilization of all detections. Conversely, a threshold of 100% signifies the inclusion of solely those detections exhibiting 100% certainty.

With increasing the threshold it can be observed that precision increases, indicating that false positives are correctly being eliminated. As for the recall, it remains almost steady until threshold of 60% and begins decreasing. This reflects the elimination of true positives after this value hence reducing the recall performance of the DNN.

For that, taking the threshold at 60% gives a good trade-off between improved precision whilst maintaining recall performance. A higher threshold would risk omitting correct detections, and a lower threshold would keep unneeded false positives.

In Table 3, a comparison between before and after thresholding for average precision and recall for overall detections and detections upto 50 meters are displayed. We can see a significant improvement for the average precision over all other approaches and also for our method without thresholding. As for recall, we can see an increase over the baseline and state-of-the-art approaches, but slightly less than the Dirichlet MLE + ILVI without thresholding.

This improvement of precision stems from the fact that the DNN is able to associate high certainty to true positives and low certainty to false positives, hence achieving an efficient separation between both categories, as shown in Table 2 and Figure 4.

5 Conclusion

This paper presents a novel approach to enhance instance segmentation performance and uncertainty estimation in deep neural networks for automated driving applications. The proposed architecture combines the Dirichlet Maximum Likelihood Estimation approach, Intermediate Layer Variational Inference, and uncertainty-based thresholding to achieve more accurate instance segmentation and reliable uncertainty estimates. This integration introduces a new perspective on countering ghost detections through uncertainty thresholding. By thresholding the uncertainty, we observe a boost in the precision performance of the DNN whilst maintaining high recall performance. The demonstrated performance improvements establish the potential of this methodology for advancing the capabilities of AI-driven systems in real-world applications. Future research can further explore the optimization of threshold selection to enhance the reliability and effectiveness of uncertainty estimation in DNNs for automated driving applications.

ACKNOWLEDGMENT

This work is partly funded by the German Federal Ministry for Economic Affairs and Climate Action (BMWK) and partly financed by the European Union in the frame of NextGenerationEU within the project "Solutions and Technologies for Automated Driving in Town" (FKZ 19A22006P).

References

- [1] Alexander Selvikvåg Lundervold and Arvid Lundervold. An overview of deep learning in medical imaging focusing on mri. *Zeitschrift für Medizinische Physik*, 29(2):102–127, 2019.
- [2] Ekim Yurtsever, Jacob Lambert, Alexander Carballo, and Kazuya Takeda. A survey of autonomous driving: Common practices and emerging technologies. *IEEE access*, 8:58443–58469, 2020.
- [3] Keli Huang, Botian Shi, Xiang Li, Xin Li, Siyuan Huang, and Yikang Li. Multi-modal sensor fusion for auto driving perception: A survey. *arXiv preprint arXiv:2202.02703*, 2022.
- [4] Szilárd Aradi. Survey of deep reinforcement learning for motion planning of autonomous vehicles. *IEEE Transactions on Intelligent Transportation Systems*, 2020.
- [5] Jingdong Wang, Ke Sun, Tianheng Cheng, Borui Jiang, Chaorui Deng, Yang Zhao, Dong Liu, Yadong Mu, Minghui Tan, Xinggang Wang, et al. Deep high-resolution representation learning for visual recognition. *IEEE transactions on pattern analysis and machine intelligence*, 2020.
- [6] Chuan Guo, Geoff Pleiss, Yu Sun, and Kilian Q Weinberger. On calibration of modern neural networks. In *International conference on machine learning*, pages 1321–1330. PMLR, 2017.

- [7] Oliver Willers, Sebastian Sudholt, Shervin Raafatnia, and Stephanie Abrecht. Safety concerns and mitigation approaches regarding the use of deep learning in safety-critical perception tasks. In *International Conference on Computer Safety, Reliability, and Security*, pages 336–350. Springer, 2020.
- [8] Alex Kendall. *Geometry and uncertainty in deep learning for computer vision*. PhD thesis, University of Cambridge, UK, 2019.
- [9] Yarin Gal. Uncertainty in deep learning. 2016.
- [10] Moloud Abdar, Farhad Pourpanah, Sadiq Hussain, Dana Rezazadegan, Li Liu, Mohammad Ghavamzadeh, Paul Fieguth, Xiaochun Cao, Abbas Khosravi, U Rajendra Acharya, et al. A review of uncertainty quantification in deep learning: Techniques, applications and challenges. *Information Fusion*, 2021.
- [11] Ahmed Hammam, Seyed Eghbal Ghobadi, Frank Bonarens, and Christoph Stiller. Real-time uncertainty estimation based on intermediate layer variational inference. In *Computer Science in Cars Symposium*, pages 1–9, 2021.
- [12] Ahmed Hammam, Frank Bonarens, Seyed Eghbal Ghobadi, and Christoph Stiller. Predictive uncertainty quantification of deep neural networks using dirichlet distributions. In *Proceedings of the 6th ACM Computer Science in Cars Symposium*, pages 1–10, 2022.
- [13] Ahmed Hammam, Frank Bonarens, Seyed Eghbal Ghobadi, and Christoph Stiller. Towards improved intermediate layer variational inference for uncertainty estimation. In *European Conference on Computer Vision*, pages 1–14. Springer, 2022.
- [14] Jakob Gawlikowski, Cedrique Rovile Njieutcheu Tassi, Mohsin Ali, Jongseok Lee, Matthias Hünt, Jianxiang Feng, Anna Kruspe, Rudolph Triebel, Peter Jung, Ribana Roscher, et al. A survey of uncertainty in deep neural networks. *arXiv preprint arXiv:2107.03342*, 2021.
- [15] Hao Wang and Dit-Yan Yeung. A survey on bayesian deep learning. *ACM Computing Surveys (CSUR)*, 53(5):1–37, 2020.
- [16] Andrey Malinin and Mark Gales. Predictive uncertainty estimation via prior networks. *Advances in neural information processing systems*, 31, 2018.
- [17] Murat Sensoy, Lance Kaplan, and Melih Kandemir. Evidential deep learning to quantify classification uncertainty. *Advances in neural information processing systems*, 31, 2018.
- [18] Arthur P Dempster. Upper and lower probabilities induced by a multivalued mapping. In *Classic works of the Dempster-Shafer theory of belief functions*, pages 57–72. Springer, 2008.
- [19] Andrew Howard, Mark Sandler, Grace Chu, Liang-Chieh Chen, Bo Chen, Mingxing Tan, Weijun Wang, Yukun Zhu, Ruoming Pang, Vijay Vasudevan, Quoc V. Le, and Hartwig Adam. Searching for mobilenetv3. In *Proceedings of the IEEE/CVF*

- International Conference on Computer Vision (ICCV)*, pages 1314–1324, October 2019.
- [20] Bowen Cheng, Maxwell D Collins, Yukun Zhu, Ting Liu, Thomas S Huang, Hartwig Adam, and Liang-Chieh Chen. Panoptic-deeplab: A simple, strong, and fast baseline for bottom-up panoptic segmentation. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 12475–12485, 2020.
 - [21] Ahmed Hammam, Frank Bonarens, Seyed Eghbal Ghobadi, and Christoph Stiller. Identifying out-of-domain objects with dirichlet deep neural networks. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) Workshops*, October 2023.
 - [22] Thomas Minka. Estimating a dirichlet distribution, 2000.
 - [23] Durk P Kingma, Tim Salimans, and Max Welling. Variational dropout and the local reparameterization trick. *Advances in neural information processing systems*, 28, 2015.
 - [24] Marius Cordts, Mohamed Omran, Sebastian Ramos, Timo Rehfeld, Markus Enzweiler, Rodrigo Benenson, Uwe Franke, Stefan Roth, and Bernt Schiele. The cityscapes dataset for semantic urban scene understanding. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 3213–3223, 2016.
 - [25] Andreas Geiger, Philip Lenz, Christoph Stiller, and Raquel Urtasun. Vision meets robotics: The kitti dataset. *The International Journal of Robotics Research*, 32(11):1231–1237, 2013.
 - [26] Jishnu Mukhoti and Yarin Gal. Evaluating bayesian deep learning methods for semantic segmentation. *arXiv preprint arXiv:1811.12709*, 2018.

Transfer Learning Techniques Using Simulation Data For Machine Learning Automotive Radar Systems

Felix Rutz^{*}, Ralph Rasshofer[†] und Erwin Biebl[‡]

Abstract: For a reliable detection and classification of vulnerable road users in modern automotive radar systems, the latest research introduces machine-learning (ML) based algorithms. However, suitable training datasets for ML systems based on real-world radar measurements are rarely available or lack specific raw radar data. Different approaches based on transfer-learning methods from data generated by a simulation framework for the range-Doppler-representation of radar measurement data are researched. In particular, influences of dataset size and sample quality, as well as different transfer learning approaches concerning the performance of the ML system in the radar data domain, are examined.

Index Terms: Automotive Radar, Machine Learning, Radar Dataset Simulation, Transfer Learning

1 Introduction

The EU Vision Zero road traffic safety initiative seeks to decrease road injuries and fatalities by 2050 [1]. Therefore, advanced driver-assistance technologies are required in modern automobiles to achieve this goal. The foundation of these systems to perform successfully is a thorough sensing of the vehicle's surroundings. Hence, radar and lidar scanners, front cameras, and ultrasonic probes comprise a basic sensor ensemble enabling a multimodal data-driven depiction of the surrounding automobile environment. The sensor-specific information obtained is subsequently analyzed and optimized for different safety-related applications, such as autonomous emergency braking, blind spot detection, and higher automated driving capabilities [2].

With the implementation of advanced conditional and automated driving functions, the various raw sensor data are often combined and evaluated by elaborately trained machine learning systems. Their performance levels are unprecedented compared to standard signal-processing methods. For example, a deep learning system is used as an evaluation framework joining the vision and the motion planning domain for improved pedestrian detection [3]. However, the application of artificial intelligence-based functions is not limited to safety-related systems. In [4], a deep learning system regulates a complex automatic torque converter transmission, outperforming the control performance of classical control approaches.

^{*}Professur für Höchstfrequenztechnik, Technische Universität München (e-mail: felix.rutz@tum.de).

[†]BMW AG, München (e-mail: Ralph.Rasshofer@bmw.de).

[‡]Professur für Höchstfrequenztechnik, Technische Universität München (e-mail: biebl@tum.de).

2 Related Work

For improved protection of vulnerable road users, radar systems, camera devices, and lidar sensors have been investigated for reliable identification and categorization of target objects. Comparing these available measurement devices, the advantage of radar sensors is their least-imperishable characteristic in changing weather or lightning conditions, even on long detection ranges, as shown by [5]. However, the advantage is not limited to axial perception, as component supplier Bosch lately introduced synthetic aperture functionality to the vehicle for a detailed lateral perception [6].

The latest generation of commercially available vehicular long-range radars uses continuous frequency-modulated ramps for target detection in the 76 – 77 GHz band [7]. This frequency band allows the use of broadband waveforms, which have advantageous effects on range and velocity resolutions. Further, the angular resolution is enhanced by antenna arrays due to the small mechanical dimensions corresponding with the required wavelength.

In order to meet the requirements of self-driving vehicles in terms of radar sensor characteristics, scaling the radar parameters, for example, by higher bandwidths, is not sufficient. In order to meet the demands of environment sensing, fundamentally different approaches are required. For this purpose, research is being conducted in the field of signal processing on more complex algorithms. Approaches include the evaluation of large array structures using sophisticated models for improved angular resolution [8], as well as the use of alternative radar modulation techniques [9].

The rising capabilities of signal processing methods based on machine learning (ML) is another promising approach, further enabling the combined detection and classification of remote objects rather than only detecting the presence of an unclassified canonical target. In [10] the range-Doppler-representation of radar measurements is used for detection tasks, whereas [11] presents a detection algorithm using the range-azimuth spectrum instead.

Common to machine-learning approaches is the limited availability of sufficiently large training datasets. Hence, the presented deep neural network algorithms are often tailored to work well on a specific input dataset limited to the situational events represented within the learning set. Therefore, a direct comparison of different approaches remains impossible, as datasets and networks are incompatible with each other [12].

For a reliable inference and generalization to unseen driving scenarios, a sufficient training database with ideally all possible situations as data samples is required for all ML-based approaches. As such datasets are generally unavailable, learning from computer-generated radar data is under research. We, therefore, create simulation datasets and use them to train a computer vision object detection system based on the YOLOv5 framework [13] on the radar data domain. Different training experiments focus on the effects of using various simulated dataset sizes. The results are evaluated for sample quantity and quality of the dataset. The corresponding ML models are then separately re-trained with input data based on radar sensor measurements using different transfer learning approaches to improve the system predictions.

Rather than solely adapting the ML model to a new domain by freezing specific layers as feature representations and re-training them with different datasets as depicted in [14], we also introduce modifications to the feature space so the model adapts well to a slightly different domain within the same training run.

3 Evaluation Metrics

For the evaluation of the machine learning system, the precision, recall, and mean average precision metrics are used following the remarks by [15]. The precision is defined as the ratio of true positives (TP) and the total number of predicted positives as a sum of the true positives and false positives (FP):

$$Precision = \frac{TP}{TP + FP}. \quad (1)$$

The precision metric expresses the accuracy of the neural network as the proportion of its correct predictions to all positive predictions. Recall or sensitivity is defined as the ratio of true positives and the total of ground truth positives as a sum of the true positives and false negatives (FN):

$$Recall = \frac{TP}{TP + FN}. \quad (2)$$

The recall metric is related to the ability of the machine learning model to find all the positive samples from the dataset. The classification system is characterized by merging both parameters into the precision over recall curve. The area under the curve is summarized as a single number used as a numeric performance indicator referred to as the average precision (AP) for each class. Deriving the mean of the AP over all distinct curves for each class leads to the mean average precision (mAP) as a general evaluation metric of an object detection system.

4 Dataset Generation from Simulation

The ML framework model performance is significantly determined by data quality and quantity used for supervised training. For the ML model examinations, we use two separate datasets. One dataset was derived from measurements with the radarbook experimental platform. The device provides a millimeter-wave radar measurement setup with raw data processing capabilities in the 76 – 77 GHz band, incorporating similar properties in frequency, modulation schemes, and bandwidth as standard automotive-grade radar sensors. The semi-manual labeling process is described in detail by [16].

The primary population of the processed measurement samples is split into a training and a test set for evaluating the ML model on unseen data. The training set is derived from 2445 pedestrian labels, 1527 bicycle annotations, 785 automobile boxes, and 562 frames with empty boxes. The test set contains 204 pedestrian samples, 553 bicycle annotations, and 95 automobile labels. The dataset is assumed to be relatively small. Regarding class distribution, the allocations of the target classes within the distinct subsets are pretty unbalanced.

Nevertheless, processing real-world sensor data and deriving large and balanced datasets is an elaborative and cost-intensive task. Identifying and labeling rare but dangerous driving scenes for their representation in the training dataset is a nearly unachievable undertaking, often also accompanied by legal implications.

Instead of manually labeling more data from an actual radar sensor, a simulated dataset is created based on the MATLAB simulation framework introduced by [17]. The cyclist, vehicle, and pedestrian objects are modeled as a sum of their characteristic reflection points. Each point target is assigned a unique position, velocity, and associated characteristic backscattering cross-section. Based on the different reflection models for pedestrians, bicyclists, and vehicles, the simulation framework calculates a radar time-domain baseband signal using the radar range formula. Multi-path propagation is excluded in order to reduce the complexity of the simulation.

The simulated sensor data is then translated into range-Doppler-map representations with appropriate ground-truth annotation labels. The synthetically generated dataset has three times as many samples as the sensor-based dataset derived from measurements, with 18 300 sample images.

A total of 12 200 samples from simulated data are used for training. The training example set contains multi-labeled samples with 8311 pedestrian labels, 10 309 bicycle annotations, and 9230 automobile boxes. The test set consists of the remaining 6100 samples. The sample distribution per class is more balanced within the simulated radar dataset than the measurement-based annotations. Table 1 compares the total number of range-Doppler-map samples and the total number of class labels for the three class entities, car, bicycle, and person, between the original sensor dataset and the simulated one.

Table 1: Comparison of manually labeled and simulated datasets

	Radar Sensor Dataset	Simulation Dataset
Total sample count	4889	18 300
No. of person labels	2649	13 219
No. of bicycle labels	2080	16 226
No. of car labels	880	12 946

5 Analysis of Simulation Data

The ability of a machine learning model to generalize and fit a dataset strongly correlates with the dataset size and data quality itself. Based on the Bagging method [18], N multiple training sets from the 12 200 simulation base population X have been drawn. However, the samples were drawn without replacement for data quantity and quality analysis concerning the ML model prediction, but with increasing the total subset size N in a range from 500 to 12 000. The step size between two dataset iterations is 500 additional samples. This prevents possibly faulty samples from being drawn into a training set multiple times.

For every subset, a separate ML model was trained from scratch. Fig. 1 depicts the results of the corresponding mAP of the resulting ML models in dependence on the different training datasets.

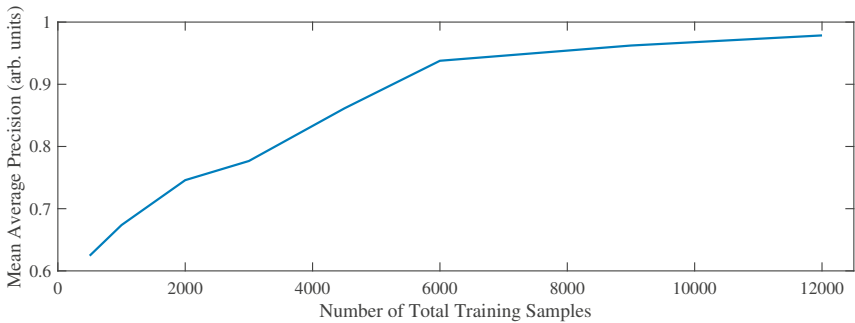
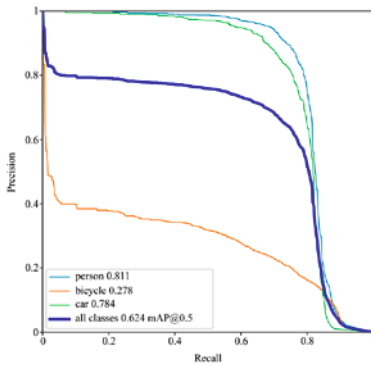


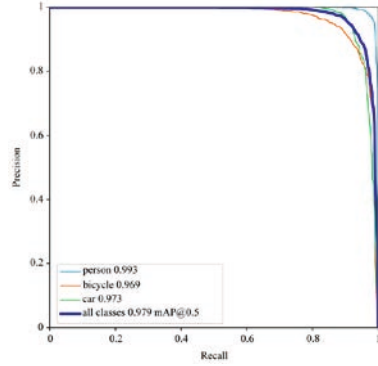
Figure 1: Evaluation of mean Average Precision with the total dataset sample size based on simulated data

On the smallest available dataset containing only 500 training samples, the ML system reaches a mean average precision of 0.62. The metric increases steeply with every increase of the training dataset until it reaches a pivot point in mAP of 0.95 using a training dataset of 6000 samples. Further experiments only correlate with minor improvements in the overall mAP score when increasing the training database size, reaching a mAP limit of 0.97 using 12 000 training examples.

The corresponding parameter models for the before-mentioned supporting nodes using 500, 6000, and 12 000 training samples are selected for further experiments. The mean average precision shows a high deviation at both system boundaries. Therefore, the detailed precision-recall plots for each class for the smallest and largest datasets are examined. Deriving from the definition of precision and recall, the combined plot summarises the trade-off between the true positive rate and the positive predictive value for an ML model [19]. Fig. 2 depicts both plots for training datasets containing 500 and 12 000 total samples.



(a) Precision-Recall-Curve for 500 samples



(b) Precision-Recall-Curve for 12000 samples

Figure 2: Comparison of Precision-Recall-Curves for training sets including 500 and 12000 total samples based on simulation data

In a qualitative comparison of both graphics, the smaller dataset shows significant setbacks in the total area covered by the precision over recall graphs for all three classes. The most noticeable flaw is discovered for the bicycle class, resulting in a mAP value of 0.278 using 500 training samples. Contrarily, the area under the curve reaches its maximum when training with 12000 samples, leading to a nearly perfect precision over recall for all classes.

Therefore, the dataset for the small-sized training set has been further investigated. Every sample image from the small dataset was inspected by hand. In the total number of bicycle instances, corresponding bounding boxes for 14 samples have been identified as misplaced. Hence, the samples only contain background noise instead of the true target ground truth representation. With an increasing number of training samples, the influence of misplaced labels seems to be mitigated within the ML model. The reason for the misplaced bounding boxes remains unclear. A likely explanation is an error during the transformation of ground truth labels into a corresponding data format, as different ML algorithms require different representations of ground truth labeling parameters.

6 Transfer to Measurement Data Domain

To implement an ML-based system in a vehicle, the models generated from simulation data must also generalize in sensor measurement data. The underlying system parameter weights from simulation data are therefore used to improve the generalization capabilities in a slightly changed domain setting. The system needs to adapt to the environmental radar data domain rather than purely classify data based on calculations excluding real-world effects such as multi-path propagation. Assuming that many factors and model weights leading to the results in simulation data also apply to the radar sensor data domain, the adaptation process of the ML model is referred to as transfer learning [20].

In traditional machine learning, a specifically tailored algorithm is trained and implemented for one detailed task using a curated dataset representing the problem. However, this approach requires the unseen input data to share the exact distribution and feature space with the available training data. This requirement is not fulfilled in most real-world applications, as shown in the previous section comparing the sample distribution in the range-Doppler-map datasets. Using transfer learning (TL) instead, the ML system can re-use its knowledge and skills of a specific task to solve a problem in another target domain. Therefore, the need for high-quality data in the target domain is reduced.

6.1 Transfer Learning by Updating Feature Representation

Following the explanations in [21], scalar feature weights represent knowledge in a machine-learning model in a network of processing units. Each processing element implements a nonlinear function, altering its input data with a specific weight. Multiple units form a specific network structure, often consisting of subsequent layers. For transfer learning, as many weights as possible are re-used from training in the source domain and fine-tuned with a small amount of data in the target domain. During the target training process, the element weights are updated, enabling the model to acquire more information to represent the input data. Different weights need to be evolved depending on the chosen ML system architecture.

Refined in detail by [22], two key aspects affect the outcome of feature adaptation, one factor being the size of the target dataset. If the target dataset is small, overfitting in the network is avoided by freezing as many layers as possible. Hence, the ML model relies more on features extracted from the source set samples. The second factor is the similarity between the source and the target dataset. With limited variations in the dataset samples, faster fine-tuning results are achievable. Finding the optimal number of layers that need to be updated or fixed during training is an incidental challenge in transfer learning. From the presented model structure, different configurations of layers are being evolved during transfer training runs with the target sensor dataset researched.

6.2 Transfer Learning by Feature Space Modification

Another approach for domain adaptation is presented in [23]. Instead of changing the network weights using multiple training runs on various source and target datasets, the sample set is expanded. By simply merging samples from both database populations into a joint training dataset, the adjacent learning process by fitting an ML model from scratch is compelled to derive a more general network representation, as samples from both data domains are considered during the same training session. Due to the continuous weight updates within each training epoch, the ML system focuses on identifying underlying features common to all input samples. Hence, a suitable generalization of the final model is achieved, and simultaneous overall training time is reduced. This straightforward approach is implemented and compared to the procedure based on updating the feature representation using two distinct datasets.

7 Transfer Learning Results

With the ML model based on 12 000 simulation samples, different transfer learning techniques have been applied using the radar sensor data. The best-considered results from the different approaches are compared with traditional ML training using only the sensor dataset as a default baseline model. Table 2 summarizes the results of the distinct transfer learning methods described before. For bicycle and car classes best results are achieved by training the ML model from scratch using a mixed dataset containing simulated and real-world data. Nonetheless, transfer learning using distinct datasets yields the best average precision for person class detection.

Table 2: Results of different transfer learning techniques

ML Training Approach	AP (Person)	AP (Bicycle)	AP (Car)	mAP@0.5
Trad. ML on Sensor Data	0.690	0.670	0.600	0.653
TL (Feature Rep.)	0.730	0.696	0.541	0.656
TL (Feature Space)	0.687	0.713	0.719	0.706

8 Conclusion

Different transfer learning techniques have been researched to improve ML model generalization in automotive radar domain data. The additional parameter adaptation from simulated radar domain data can mitigate the limitations of small sensor datasets for training. The simulation framework generates scalable datasets, including the corresponding ground-truth labels, for pre-training quickly and reliably.

However, the decrease in mAP when transferring the model weights from the simulation to the measurement data domain indicates that more underlying effective mechanisms exist in the sensor radar domain. Nevertheless, transfer learning methods significantly increase the algorithm’s performance compared to the default baseline model derived from the traditional system training approach.

Depending on the distinct object class, precision is increased up to 11.9 percent using transfer learning approaches for model generalization. The requirement for refining datasets from raw sensor data containing an even distribution of common and rare events is reduced by adding enormous samples from simulation frameworks.

Future research focuses on further improvements in generalization by data augmentation of sensor-based datasets by upsampling the total amount of available training pairs, which is directly related to further reducing data sample collection efforts.

References

- [1] European Climate, Infrastructure and Environment Executive Agency (CINEA), *EU Road Safety: Towards "Vision Zero"*, European Union, 2022.
- [2] H. Winner, S. Hakuli, F. Lotz, and C. Singer, Eds., *Handbuch Fahrerassistenzsysteme: Grundlagen, Komponenten und Systeme für aktive Sicherheit und Komfort*, 3rd ed., ser. ATZ-MTZ-Fachbuch. Wiesbaden: Springer Vieweg, 2015.
- [3] M. Lyssenko, C. Gladisch, C. Heinzemann, M. Woehle, and R. Triebel, "Towards safety-aware pedestrian detection in autonomous systems," in *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2022, pp. 293–300.
- [4] G. Gaiselmann, S. Altenburg, S. Studer, and S. Peters, "Deep Reinforcement Learning for Gearshift Controllers in Automatic Transmissions," *Array*, vol. 15, p. 100235, 2022.
- [5] A. Mukhtar, L. Xia, and T. B. Tang, "Vehicle Detection Techniques for Collision Avoidance Systems: A Review," *IEEE Transactions on Intelligent Transportation Systems*, vol. 16, no. 5, pp. 2318–2338, 2015.
- [6] M. Farhadi, R. Feger, J. Fink, T. Wagner, and A. Stelzer, "Automotive synthetic aperture radar imaging using tdm-mimo," in *2021 IEEE Radar Conference (Radar-Conf21)*, 2021, pp. 1–6.
- [7] K. Ramasubramanian and K. Ramaiah, "Moving from Legacy 24 GHz to State-of-the-Art 77-GHz Radar," *ATZelektronik worldwide*, vol. 13, no. 3, pp. 46–49, 2018. [Online]. Available: <https://doi.org/10.1007/s38314-018-0029-6>
- [8] G. Schnoering, C. Höller, T. Kawaguchi, K. Kawajiri, and S. Malterer, "Fast angular processing for sparse fmcw radar arrays with non-uniform fft," in *2023 24th International Radar Symposium (IRS)*, 2023, pp. 1–11.
- [9] G. Hakobyan and B. Yang, "High-Performance Automotive Radar: A Review of Signal Processing Algorithms and Modulation Schemes," *IEEE Signal Processing Magazine*, vol. 36, no. 5, pp. 32–44, 2019.
- [10] W. Ng, G. Wang, Siddhartha, Z. Lin, and B. J. Dutta, "Range-doppler detection in automotive radar with deep learning," in *2020 International Joint Conference on Neural Networks (IJCNN)*, 2020, pp. 1–8.
- [11] K. Patel, K. Rambach, T. Visentin, D. Rusev, M. Pfeiffer, and B. Yang, "Deep learning-based object classification on automotive radar spectra," in *2019 IEEE Radar Conference (RadarConf)*, 2019, pp. 1–6.
- [12] D. Feng, C. Haase-Schütz, L. Rosenbaum, H. Hertlein, C. Gläser, F. Timm, W. Wiesbeck, and K. Dietmayer, "Deep Multi-Modal Object Detection and Semantic Segmentation for Autonomous Driving: Datasets, Methods, and Challenges," *IEEE Transactions on Intelligent Transportation Systems*, vol. 22, no. 3, pp. 1341–1360, 2021.

- [13] G. Jocher *et al.*, “ultralytics/yolov5: v6.2 - YOLOv5 Classification Models, Apple M1, Reproducibility, ClearML and Deci.ai integrations,” Aug. 2022. [Online]. Available: <https://doi.org/10.5281/zenodo.7002879>
- [14] F. Rutz, E. Biebl, and J.-P. Konrad, “Transfer Learning in ML-based Radar Systems for Automotive Applications,” in *2022 Kleinheubach Conference*, 2022, pp. 1–3.
- [15] K. P. Murphy, *Machine learning: A probabilistic perspective*, 4th ed., ser. Adaptive computation and machine learning series. Cambridge, Mass.: MIT Press, 2013.
- [16] R. Pérez, F. Schubert, R. Rasshofer, and E. Biebl, “Deep learning radar object detection and classification for urban automotive scenarios,” in *2019 Kleinheubach Conference*, 2019, pp. 1–4.
- [17] T. Wengerter, R. Pérez, E. Biebl, J. Worms, and D. O’Hagan, “Simulation of urban automotive radar measurements for deep learning target detection,” in *2022 IEEE Intelligent Vehicles Symposium (IV)*, 2022, pp. 309–314.
- [18] E. Alpaydin, *Introduction to Machine Learning*, 4th ed., ser. Adaptive Computation and Machine Learning. Cambridge, Massachusetts: The MIT Press, 2020.
- [19] T. Saito and M. Rehmsmeier, “The precision-recall plot is more informative than the roc plot when evaluating binary classifiers on imbalanced datasets,” *PloS one*, vol. 10, no. 3, p. e0118432, 2015.
- [20] I. Goodfellow, Y. Bengio, and A. Courville, *Deep learning*, ser. Adaptive Computation and Machine Learning. Cambridge, Massachusetts: The MIT Press, 2016.
- [21] R. Reed and R. J. Marks, *Neural Smithing: Supervised Learning in Feedforward Artificial Neural Networks*. The MIT Press, 02 1999. [Online]. Available: <https://doi.org/10.7551/mitpress/4937.001.0001>
- [22] M. Elgendy, *Deep learning for vision systems*. Shelter Island: Manning, 2020.
- [23] H. Daumé, “Frustratingly Easy Domain Adaptation,” 2009. [Online]. Available: <https://arxiv.org/abs/0907.1815>

Spieltheoretischer prädiktiver Regler für Interaktion im gemischten Verkehr

Mohamed-Khalil Bouzidi* und Ehsan Hashemi†

Zusammenfassung: Dieser Artikel präsentiert eine integrierte Trajektorienplanung und Regelung für Verkehrsszenarien in denen Interaktion mit Fahrzeugen, die von Menschen gefahren werden, erforderlich ist. Für eine sichere Bewegungsplanung berücksichtigt die vorgeschlagene Methode die Interaktion zwischen dem automatisierten Fahrzeug und anderen Fahrzeugen anhand der Spieltheorie. Das Framework umfasst ein neuartiges inverses Differentialspiel basierend auf einer LSTM Architektur, um die Zielfunktion des menschlichen Fahrers online zu schätzen. Diese Schätzung wird dann von einem spieltheoretischen prädiktiven Regler genutzt, um die Trajektorie des menschlich gesteuerten Fahrzeugs zu präzisieren und das autonome Fahrzeug zu steuern. Das entwickelte System wird in einer Human-in-the-Loop Probandenstudie unter Verwendung der CarSim High-Fidelity-Simulation evaluiert.

Schlüsselwörter: Long Short-Term Memory (LSTM), Mensch-Maschine-Interaktion, Spieltheoretischer Modellprädiktiver Regler, Trajektorienplanung

1 Einleitung

Testfahrten autonomer Fahrzeuge(ADS) haben in den letzten Jahren gezeigt, dass diese Systeme bei bestimmten Herausforderungen an ihre Grenzen stoßen. Eines dieser Probleme, das so genannte Frozen Robot Problem(FRP) [1], wurde bereits mehrfach beobachtet. Insbesondere in komplexen Verkehrsszenarien, wie z.B. bei Einfädelungsmanövern, scheitern autonome Fahrzeuge bei der Durchführung. Dieses Fehlverhalten lässt sich durch die Entkopplung der Prädiktionsaufgabe (die die Trajektorien der Verkehrsteilnehmer in der Szene prädiziert) von der Planungsaufgabe (die eine sichere Trajektorie auf der Basis einer der vorgeschalteten Prädiktionen generiert) erklären: Die umliegenden Fahrzeuge werden als dynamische Hindernisse betrachtet, die nicht auf die gefahrene Trajektorie des ADS reagieren, wodurch durchführbare Manöver nicht erkannt werden. Um in der Lage zu sein, die Situation korrekt zu antizipieren, muss das ADS wissen, dass die umliegenden Fahrzeuge mit dem ADS interagieren, um Kollisionen zu vermeiden. Dies erfordert gekoppelte Vorhersage- und Planungsaufgaben, da die Fahrer ihr Verhalten an die Trajektorien der anderen Fahrzeuge anpassen und umgekehrt.

In dieser Hinsicht bietet die Spieltheorie einen geeigneten Rahmen, um diese gleichzeitige Planung und Prädiktion anzugehen. Da im gemischten Verkehr oft Verkehrsteilnehmer nicht unbedingt kooperativ agieren, werden solche Interaktionen oft als unkooperatives Spiel z.B. als

*Mohamed-Khalil Bouzidi ist bei Continental AG tätig und Doktorand an der Freien Universität Berlin, Deutschland. Die Arbeit wurde im Rahmen seiner Masterarbeit am NODE Lab, Fachbereich Maschinenbau, University of Alberta, Edmonton, AB, Kanada, durchgeführt. (e-mail: m.bouzidi@fu-berlin.de).

†Ehsan Hashemi ist im Fachbereich Maschinenbau, University of Alberta, Edmonton, AB, Kanada, (e-mail: ehashemi@ualberta.ca).

Nash-Equilibrium (NE) modelliert. Existierende Ansätze für Trajektorienplanung [2,3] beruhen jedoch auf der Annahme, dass die Zielfunktion des Fahrers dem ADS bekannt ist.

Die Fahrer haben in der Regel unterschiedliche Ziele und Fahrstile, was eine Online-Schätzung der Zielfunktion des Menschen erforderlich macht. Zur Schätzung von Zielfunktionen wird in der Literatur das inverse Differentialspiel (IDG) verwendet. Bestehende Methoden sind jedoch entweder zu rechenintensiv, um sie online auszuführen, wie z. B. direkte Methoden [4], oder inverse Reinforcement Learning basierte Methoden [5] oder sie beruhen auf zu starken Annahmen, wie Methoden, die von Inverse Optimal Control [6] abgeleitet sind.

In [7] wird ein Framework vorgeschlagen, die ein Feed-Forward-Netz für die Online-IDG verwendet, und ein weiteres Netz, das im nächsten Schritt eine NE findet. Das Fehlen eines Fahrzeugmodells und die begrenzten (weniger als 50) Trainingsdaten bieten jedoch keine Garantie für fahrbare Trajektorien. In [4] wird ein unscented Kalman-Filter (UKF) zur Online-Schätzung der Zielfunktionsparameter verwendet (d. h. eine direkte Methode für IDG) welches dem spieltheoretische modellprädiktive Regler (GT MPC) [3] als Input dient. Die Methoden verwendet das kinematische Einspurmodell, das das Fahrzeugverhalten bei Autobahn-Einfädelmanövern, die in der Regel mit hoher Geschwindigkeit durchgeführt werden und eine ungleichmäßige Drehmomentverteilung an jeder Achse aufweisen können, nicht angemessen approximieren kann. Dies kann zur Generierung unrealisierbarer Trajektorien führen. Ein einfaches Austauschen des Modells ist allerdings nicht möglich, da eine höhere Anzahl von Zuständen und Parametern der Zielfunktion (aufgrund komplexerer dynamischer Modelle) zu rechnerischen Herausforderungen für den vorgeschlagenen UKF-basierten Schätzer führt.

Um diese Probleme anzugehen, schlagen wir ein interaktionsbewusstes Framework basierend auf der Spieltheorie vor, der das dynamische Einspurmodell verwendet, um dynamisch realisierbare Trajektorien zu generieren, und ein LSTM-Netzwerk für die Online-Schätzung des Verhaltens von Fahrzeugen bei Einfädelungen in Echtzeit verwendet. Unser Framework wurde in Einfädelungsszenarien im gemischten Verkehr evaluiert (mit menschlichen Probanden in einem Human-in-the-Loop-Setup und unter Verwendung hochgenauer *CarSim*-Simulationen), was aufgrund der zuvor genannten Gründe für andere spieltheoretische Ansätze [2, 4, 7] nicht möglich ist. Unsere Hauptbeiträge lassen sich daher wie folgt zusammenfassen:

- Entwicklung eines LSTM-basierten inversen Differentialspiels zur Online-Schätzung der Zielfunktion des menschlichen Fahrers für eine prädiktive Regelungsstrategie.
- Entwurf eines Unknown Input Observer zur Schätzung der Steuergrößen des umgebenden Fahrzeugs mithilfe von Bord-Sensorikmessungen.
- Entwicklung einer spieltheoretischen modellprädiktiven Regelung, die die laterale Fahrzeugdynamik berücksichtigt, um laterale Stabilität und dynamische Realisierbarkeit der geplanten Trajektorien für eine sicherere Mensch-ADS-Interaktion zu gewährleisten.

2 Problem Definition und Modellbeschreibung

Betrachtet wird ein autonomes Ego-Fahrzeug welches in die Autobahn einfädeln möchte. Ein von einem Menschen gesteuertes Fahrzeug (auf der linken Seite) interagiert entsprechend mit dem ADS. Das Ziel besteht darin, eine kollisionsfreie Trajektorie zu planen, dem das Ego-Fahrzeug folgen kann. Um die interaktive Natur des Autobahn-Einfädelns zu erfassen, modellieren wir dieses Szenario als ein unkooperatives Spiel. Wir wählen das NE als Lösungskonzept,

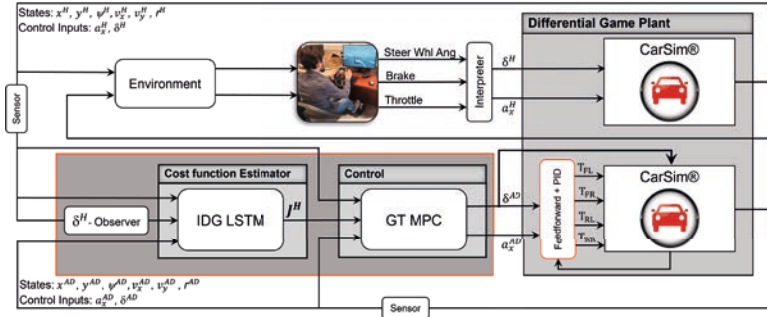


Abbildung 1: Das interaktionsbewusste spieltheoretische Trajektorienplanungs- und Regelung Framework unter Verwendung eines Unknown Input Observer und LSTM

da bei Nash-Spielen kein Spieler einen strukturellen Vorteil gegenüber anderen Spielern hat, im Gegensatz zur Leader-Follower-Struktur eines Stackelberg-Spiels.

Definition 1. Das Nash-Gleichgewicht ist ein Lösungskonzept, das entsteht, wenn jeder Spieler/Agent i gleichzeitig und optimal bezüglich seiner eigenen Zielfunktion J^i handelt und seinen Annahmen über die Strategien der anderen Spieler j , wobei diese Annahmen korrekt sind. Dies entspricht:

$$J^i(x^*, u_i^*, u_j^*) < J^i(x^*, u_i, u_j^*) \quad \forall i. \quad (1)$$

Spieltheoretisch ausgedrückt, kann dieses Autobahn-Einfädelungsszenario als eine Szene beschrieben werden, in der zwei Agenten $i \in AD, H$ (automatisiertes Fahrersystem und von einem Menschen gesteuertes Fahrzeug) in einem dynamischen Spiel mit endlichem Horizont verhandeln. Um diese Szene als Nash-Spiel zu modellieren, treffen wir folgende Annahme:

Annahme 1. Der Mensch handelt rational, um seine Zielfunktion J^H zu minimieren.

Tatsächlich wird das Entscheidungsverhalten und die Reaktion des Menschen bei Fahrmanövern oft in der Literatur [8] durch einen optimalen Regler modelliert, da Menschen das Fahrverhalten anderer basierend auf ihrem Wissen vorhersagen und darauf rational reagieren.

Es sollte beachtet werden, dass es keine Kommunikation zwischen den Fahrzeugen wie V2V gibt. Daher hat das ADS keinen Zugriff auf die Zielfunktion des anderen Fahrzeugs (d.h. Agent H). Allerdings wird diese Zielfunktion benötigt, um das NE zu berechnen. Daher wird die Zielfunktion online mittels eines IDG und der beobachteten Trajektorien der Fahrzeuge geschätzt. Mit der Schätzung der Zielfunktion wird die Lösung des NE basierend auf dem Modell des Systems berechnet, wie in Abbildung 1 veranschaulicht.

Wir modellieren das Spiel als Differentialspiel, da die Fahrzeugkinematik und -dynamik naturgemäß kontinuierlich sind. Das Differentialspiel wird durch ein dynamisches System beschrieben: $\dot{x}(t) = f(x(t), u^H(t), u^A(t))$. Aufgrund von Situationen mit hoher Schlupfgefahr beim Einfädeln verwenden wir ein dynamisches Einspurmodell, um Trajektorien präzisieren, die die Manövrierfähigkeiten (und Grenzwerte der Fahrzeugführung) korrekt einschätzen und somit dynamisch machbare Trajektorien generieren. Daher werden die seitlichen Reifenkräfte

$F_{y,f}$ und $F_{y,r}$ berücksichtigt, wobei von linearen Reifenkräften ausgegangen wird. Der Schlupfwinkel an der Vorder- und Hinterachse des Fahrzeugs wird jeweils durch $\alpha_f = \delta - \frac{l_f r + v_y}{v_x}$ und $\alpha_r = \frac{l_r r - v_y}{v_x}$ definiert. Hierbei steht δ für den Lenkwinkel der Vorderachse und die Längsgeschwindigkeit, Quergeschwindigkeit und Giergeschwindigkeit des Fahrzeugs im Fahrzeugbezugssystem (bezogen auf den Fahrzeugschwerpunkt, CG) werden durch v_x , v_y und r bezeichnet. Die Fahrzeugdynamik kann dann wie folgt beschrieben werden:

$$\dot{v}_x = v_y r + a_x, \quad \dot{v}_y = -v_x r + \left(\frac{F_{y,f} + F_{y,r}}{m} \right), \quad \dot{r} = \ddot{\psi} = \left(\frac{F_{y,f} l_f - F_{y,r} l_r}{I_z} \right) \quad (2)$$

wo a_x der Längsbeschleunigung entspricht, m ist die Fahrzeugmasse und die Abstände der Vorder- und Hinterachse zum Schwerpunkt des Fahrzeugs werden mit l_f und l_r bezeichnet. Die Dynamik der Längs-/Querposition ergibt sich zu $[\dot{x}, \dot{y}]^\top = R_\psi [v_x, v_y]^\top$ mit der Rotationsmatrix R_ψ sowohl für das ADS als auch für das von einem Menschen gesteuerte Fahrzeug. Diese Modelle werden dann gestapelt, mit Ausnahme des ersten Zustands (d.h., der Längsposition, bei der der inter-vehikuläre Abstand e_x als Zustandsvariable verwendet wird), um das dynamische Modell für das Differentialspiel mit den Steuergrößen $\mathbf{u}^i = [a_x^i, \delta^i]^\top$ für $i \in AD, H$ und den erweiterten Zuständen $\mathbf{x} = [e_x, y^{AD}, y^H, \psi^{AD}, \psi^H, v_x^{AD}, v_x^H, v_y^{AD}, v_y^H, r^{AD}, r^H]^\top$ zu bilden.

Um die Zielfunktion des menschlichen Fahrers H zu bestimmen, verwenden wir einen Basisfunktions-Ansatz. Es wird angenommen, dass die Zielfunktion beider Agenten $i \in AD, H$ durch $J^i = \int \theta^i \phi^i(t) dt$ beschrieben werden kann. Die Parameter θ^H müssen mit Hilfe eines IDG identifiziert werden. Die Parameter θ^{AD} des ADS sind einstellbar. Wir beschreiben das Ziel des Fahrzeugs H durch (3), wobei der Term zur Minimierung der Quergeschwindigkeit und somit des Seitenschlupfwinkels entscheidend ist, um die laterale Stabilität sicherzustellen.

$$J^i = \int \theta_1^i (v_x^i - v_{x,d}^i)^2 + \theta_2^i (y^i - y_d^i)^2 + \theta_3^i (v_y^i)^2 + \theta_4^i \tanh \left(\underbrace{\sqrt{\left(e_x^2 + \left(\frac{b}{a} (y^H - y^{AD}) \right)^2 \right)}}_{\Gamma} - a \right) + \theta_5^i (a_x^i)^2 + \theta_6^i (\delta^i)^2 dt \quad \forall i \quad (3)$$

Die Zielfunktion (3) kann in $J^i = \int \theta^i \phi^i(t) dt$ mit den Parametern $\theta^i = [\theta_1^i, v_{x,d}^i, \theta_2^i, y_d^i, \theta_3^i, \theta_4^i, \theta_5^i, \theta_6^i]^\top$ und der Basisfunktion $\phi^i(t) = [(v_x^i)^2, -2v_{x,d}^i, (y^i)^2, -2y_d^i, (v_y^i)^2, \tanh(\Gamma), (a_x^i)^2, (\delta^i)^2]^\top$ transformiert werden.

3 Unknown Input Observer Design

Das ADS ist mit Kameras, LiDARs und Radarsensoren zur Detektion von Verkehrsteilnehmern und zur Messung ihrer relativen Abstände/Ausrichtungen ausgestattet. Der Lenkwinkel des von einem Menschen gesteuerten Fahrzeugs ist nicht messbar. Daher wird ein Unknown Input Observer (UIO) [9] entworfen, um diesen Eingang mithilfe des linearen lateralen Fahrzeugdynamikmodells $\dot{\zeta} = A\zeta + Bu$ mit den Zuständen $\zeta = [v_y^H, r^H]^\top$, messbaren Ausgängen $\mathbf{y} = \zeta$ und dem unbekannten Eingang $\mathbf{u} = \delta^H$ abzuschätzen. Der folgende Beobachter approximiert den Lenkeingang des umgebenden Fahrzeugs, wobei $\mathbf{z}$ die Zustandsvariable für den Beobachter ist.

$$\hat{\delta}^H = [B^\top B]^{-1} B^\top (\dot{\mathbf{z}} - \hat{\zeta} - A\hat{\mathbf{y}}), \quad \dot{\mathbf{z}} = F\mathbf{z}(t) + K\mathbf{y}(t), \quad \hat{\zeta} = \mathbf{z}(t) + \bar{H}\mathbf{y}(t). \quad (4)$$

Die Fehlerdynamik des Schätzerfehler $e = \zeta - \hat{\zeta}$ entspricht dementsprechend:

$$\dot{e} = (A - \bar{H}A - K_1)e + [F - (A - \bar{H}A - K_1)]z + [K_2 - (A - \bar{H}A - K_1)]y + (\bar{H} - I)Bu, \quad (5)$$

wobei die asymptotische Stabilität der Fehlerdynamik erreicht wird, wenn die Matrizen des Beobachter wie folgt entworfen werden: $\bar{H} = B[B^\top B]^{-1}B^\top$, $F = (I - \bar{H}C)A - K_1C$, $K_2 = F\bar{H}$, wobei K_1 so gewählt wird, dass F Hurwitz ist [9]. Einfädelungsszenarien beinhalten in der Regel Geschwindigkeitsvariationen aufgrund Interaktion mit dem umgebenden Fahrzeug, weshalb die Matrizen unter Verwendung der gemessenen Geschwindigkeit adaptiert werden.

4 Spieltheoretischer Prädiktiver Regler

Unter Verwendung des nichtlinearen Systemmodells $\dot{x} = f(x, u^H, u^{AD})$ und der Zielfunktionen J^{AD} , J^H entwerfen wir ein spieltheoretisches Modellprädiktiver Regler (GT MPC), der das Open-Loop-Problem für einen gegebenen Prädiktionshorizont löst, aber die Anfangszustandswerte x_0 in jedem Zeitschritt mit dem tatsächlich beobachteten Zustand aktualisiert.

Die Lösung für das Open-Loop-NE (OLNE) kann numerisch bestimmt werden, indem es in ein nichtlineares Programm (NLP) mit den Variablen $X = [x_1, x_2, \dots, x_{N_p}]^\top$, $U^i = [u_1^i, \dots, u_{N_c}^i]^\top$ umgewandelt wird, wobei N_p den Vorhersagehorizont und N_c den Steuerungshorizont angibt. Wir formulieren das Lagrange-Multiplikator-Problem wie folgt:

$$L^i = J^i + \sum_{k=0}^{N_p} (\lambda_k^i)^\top [x_{k+1} - f_k(x_k, u_k^H, u_k^{AD})] \quad \forall i \quad (6)$$

Entsprechend müssen für ein OLNE die folgenden Optimalitätsbedingungen erfüllt sein:

$$\begin{aligned} G^i &= \nabla_{X, U^i} L^i(X, U^{AD}, U^H, \lambda^{AD}, \lambda^H) = 0 \quad \forall i \\ C &= [x_1 - f_k(x_0, u_0) \dots x_{N_p} - f_k(x_{N_p-1}, u_{N_p-1})]^\top = 0 \end{aligned} \quad (7)$$

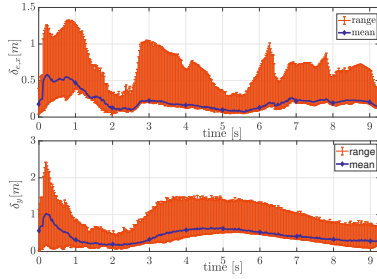
Um eine Lösung zu finden, die den Optimalitätsbedingungen genügt, verwenden wir CasADi [10] mit folgender Suchrichtung:

$$\begin{aligned} p_{ND} &= -H^{-1} G \\ G &= [G^{AD}, G^H, C]^\top, \quad H = \nabla_{(X, U, \lambda)} G(X, U^{AD}, U^H, \lambda^{AD}, \lambda^H) \end{aligned} \quad (8)$$

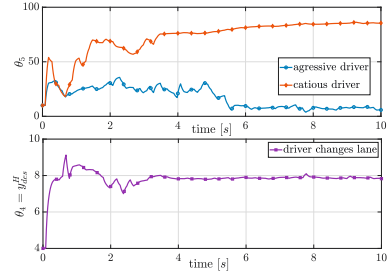
5 Inverses Differentialspiel

Die IDG dient hier dazu mithilfe von zuvor beobachteten Trajektorien $x(t)$, $u^H(t)$ und $u^{AD}(t)$ die Parameter θ^H der Zielfunktion des menschlichen Fahrers $J = \int \theta^H \phi^H(t) dt$ zu bestimmen.

Die LSTM-Struktur eignet sich gut zur Schätzung von Zuständen aus sukzessiv eintreffenden Features, wobei die Länge der Input-Zeitreihe variieren kann. Daher muss die Dauer des Manövers nicht im Voraus bekannt sein. Das Ziel besteht darin, den Parametervektor $\theta_k^H = \theta^H$ zu schätzen (d. h., ein Sequence-to-One-Regressionsproblem). Das Input Layer erhält den aktuellen Zustand x_k , die Steuergröße des menschlichen Fahrers u_k^H sowie die vorherige Steuergrößen des ADS u_{k-1}^{AD} . Dies liegt daran, dass die Entscheidung des Systems nicht vorhergesagt werden kann, indem nur das menschliche Fahrzeug isoliert beobachtet wird, da die Interaktion seine Entscheidung beeinflusst.



a Evaluation des Prädiktionsfehlers (longitudinale Distanz und laterale Position) aller Probanden



b Schätzung der Zielfunktion bei aggressiven und vorsichtigen Fahrer und bei Fahrspurwechsel

Anmerkung 1. Bei Nash-Spielen ist bekannt, dass jeder Spieler seine Basisfunktion minimieren möchte, um die Zielfunktion mit positiven Gewichten der Basisfunktionen zu minimieren. Zu jedem Zeitschritt wird die Basisfunktion $\phi_k^H t_k$ berechnet und als Input für das neuronale Netz verwendet. Dieser zusätzliche Input dient dazu Vorwissen einzubeziehen und die Generalisierungsleistung des Netzes zu verbessern.

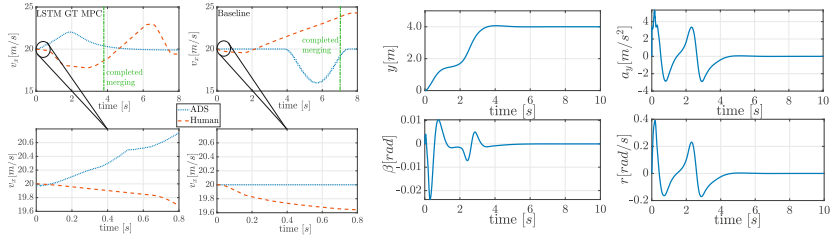
Das neuronale Netz dient als kontinuierliche Abbildung mit nichtlinearen Dynamiken mit $\theta^H = \mathcal{N}(x_k, u_k^H, u_{k-1}^{AD}, \phi_k^H)$. Es besteht aus zwei LSTM-Layer mit jeweils 80 und 60 hidden units. Dahinter werden zwei fully connected Layer mit 40 Neuronen und der RELU-Aktivierungsfunktion geschaltet.

Ein GT MPC wird verwendet, um zufällige Trajektorien synthetisch für das Training zu generieren, wobei angenommen wird, dass die von GT MPC generierten Trajektorien ausreichend konsistent mit realen Fahrzenarien unter Beteiligung von Menschen sind. Dies wurde durch umfangreiche Tests im letzten Abschnitt bestätigt. Dementsprechend werden unterschiedliche Längen von Trajektorien für das Training verwendet, so dass das LSTM-Modell den Parametervektor schätzen kann, nachdem es einen Teil der Trajektorie beobachtet hat.

6 Evaluation

Die Evaluation unseres Ansatzes wird anhand einer Probandenstudie mit 6 verschiedenen Personen durchgeführt, die nicht mit dem Algorithmus vertraut sind. Das Ziel der Studie ist es, zu analysieren, ob die spieltheoretische Modellierung für diesen Anwendungsfall gerechtfertigt ist (d.h. ob Annahme 1 gerechtfertigt ist) und ob das Framework robust gegenüber unterschiedlichen Fahrstilen ist. Für den Fahrsimulator verwenden wir highfidelity Simulationen in *CarSim*. Die Probanden nutzen zweimal das Human-in-the-Loop-Setup. Die erste Hälfte der Teilnehmer beginnt mit einem konventionellen Planer [11], die zweite Hälfte mit dem neuen Planer. Für das autonome Fahrzeug haben wir ein anspruchsvolles Szenario gewählt, bei dem beide Fahrzeuge mit derselben Anfangsgeschwindigkeit von $v_{x,0} = 20$, m/s und einem Längsabstand von $e_{x,0} = 1$, m starten.

In allen Fällen war unser Ansatz in der Lage, ein sicheres Einfädeln durchzuführen. Das Verhältnis des erfolgreichen Einfädelns von vorne betrug 4/6. Die durchschnittliche Einfädelzeit betrug 4.26 Sekunden. Der Baseline-Planer konnte das Einfädelmanöver in 5/6 der Fälle erfolgreich abschließen, und das Verhältnis des erfolgreichen Einfädelns von vorne betrug 3/6.



a Vergleich des neuen Ansatzes mit dem Baseline-Planer an einem ausgewählten Beispiel b Laterales Verhalten des Autonomen Fahrzeugs während des ausgewählten Beispiels

Die durchschnittliche Einfädelzeit betrug 7.18 Sekunden. Die durchschnittliche Einfädelzeit und das erfolgreiche Einfädeln von vorne sind wichtige Indikatoren dafür, wie erfolgreich das autonome Fahrzeug in Einfädel-Szenarien sein wird, da sich hinter dem aktuellen menschlichen Fahrzeug ein weiteres menschliches Fahrzeug befinden könnte, mit dem das autonome Fahrsystem erneut interagieren muss. Die Ergebnisse zeigen daher, dass unser Planer weniger anfällig für das Frozen Robot Problem ist.

Die Ursache hierfür lässt sich anhand eines Beispiels (siehe Abb. 3a) erklären. Hier bremst der menschliche Fahrer leicht ab, um dem autonomen Fahrzeug das Einfädeln zu signalisieren. Das entwickelte interaktionsbewusste Framework antizipiert die Reaktion des anderen Fahrers für ein sicheres Einfädeln, d.h. es kann vorhersehen, dass der Mensch leicht bremst, weil er mit dem autonomen System interagiert, und in Zukunft stärker bremsen wird. Dieses Vorwissen führt zu einer verbesserten Vorhersage der Trajektorien (siehe Abb. 2a), wobei auch zwischen aggressiven und vorsichtigen Fahrern, unter anderem anhand des Wertes der Kollisionsvermeidungsgewichtung θ_5 , unterschieden werden kann (siehe Abb. 2b). So weiß das autonome Fahrzeug früher, wann es beschleunigen (oder bremsen) sollte, und kann sicher auf die Autobahn einfädeln, wobei die laterale Stabilität trotz der Abweichung zwischen dem angenommenen Einfachspurmodell und *CarSim* aufrechterhalten wird (z.B. siehe Abb. 3b). Auf der anderen Seite wartet der Baseline-Planer in diesem Beispiel darauf, dass der Mensch stärker bremst, weil er nicht erkennt, dass der Mensch auf das autonome System reagiert. Als Ergebnis reagiert der Baseline-Regler nicht, das menschlich gesteuerte Fahrzeug beschleunigt und das autonome Fahrzeug muss hinter dem Fahrzeug einfädeln. Dieses Verhalten führt zum FFrozen Robot Problem im dichten Verkehr, wenn die nachfolgenden Fahrzeuge ähnlich handeln.

7 Fazit

Ein interaktionsbewusstes Framework für Trajektorienplanung und -regelung wurde entwickelt, der in der Lage ist, mit einem vom Menschen gesteuerten Fahrzeug zu interagieren. Es werden Einfädelungsszenarien betrachtet, welche wir als Nash-Differentialspiel mit Hilfe des dynamischen Einspurmodells beschreiben. Basierend darauf wurde ein spieltheoretische prädiktiver Regler entwickelt und getestet. Der prädiktive Regler nimmt die Zielfunktion des menschlichen Fahrers des anderen Fahrzeugs als Input, die in Echtzeit mit Hilfe eines lernbasierten inversen Differentialspiels geschätzt wird. Die experimentelle Auswertung in *CarSim* mit einem Human-

in-the-Loop-Setup hat gezeigt, dass die spieltheoretische Modellierung solcher Szenarien vorteilhaft für diese Art von Szenarien und der interaktionsbewusste Regelungssystem in der Lage ist, verschiedene Fahrstile der von Menschen gesteuerten (d.h. umgebenden) Fahrzeuge mit guter Recheneffizienz zu handhaben.

Literatur

- [1] P. Trautman and A. Krause, “Unfreezing the robot: Navigation in dense, interacting crowds,” in *2010 IEEE/RSJ International Conference on Intelligent Robots and Systems*, 2010, pp. 797–803.
- [2] D. Fridovich-Keil, E. Ratner, L. Peters, A. Dragan, and C. Tomlin, “Efficient iterative linear-quadratic approximations for nonlinear multi-player general-sum differential games,” 05 2020, pp. 1475–1481.
- [3] S. L. Cleac'h, M. Schwager, and Z. Manchester, “ALGAMES: A fast solver for constrained dynamic games,” in *Robotics: Science and Systems XVI*. Robotics: Science and Systems Foundation, 2020.
- [4] S. Le Cleac’h, M. Schwager, and Z. Manchester, “Lucidgames: Online unscented inverse dynamic games for adaptive trajectory prediction and planning,” *IEEE Robotics and Automation Letters*, vol. 6, no. 3, pp. 5485–5492, 2021.
- [5] N. Mehr, M. Wang, and M. Schwager, “Maximum-entropy multi-agent dynamic games: Forward and inverse solutions,” 2021.
- [6] J. Inga, *Inverse Dynamic Game Methods for Identification of Cooperative System Behavior*. Karlsruher Institut für Technologie, 2020.
- [7] P. Geiger and C.-N. Straehle, “Learning game-theoretic models of multiagent trajectories using implicit layers,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 35.
- [8] C. C. MacAdam, “Understanding and modeling the human driver,” *Vehicle System Dynamics*, vol. 40, pp. 101 – 134, 2003.
- [9] M. Darouach, M. Zasadzinski, and S. J. Xu, “Full-order observers for linear systems with unknown inputs,” *IEEE Transactions on Automatic Control*, vol. 39, no. 3, pp. 606–609, 1994.
- [10] J. A. E. Andersson, J. Gillis, G. Horn, J. B. Rawlings, and M. Diehl, “CasADi – A software framework for nonlinear optimization and optimal control,” *Mathematical Programming Computation*, vol. 11, no. 1, pp. 1–36, 2019.
- [11] Y. Ding, W. Zhuang, L. Wang, J. Liu, L. Guvenc, and Z. Li, “Safe and optimal lane-change path planning for automated driving,” *Proceedings of the Institution of Mechanical Engineers, Part D: Journal of Automobile Engineering*, vol. 235, no. 4, pp. 1070–1083, 2021.

Automatisiertes Kreuzungsmanagement im Urbanen Mischverkehr mit Reinforcement Learning

Marvin Klimke^{*†} Benjamin Völz^{*} Michael Buchholz[†]

Zusammenfassung: Das vernetzte automatisierte Fahren hat das Potential, die Verkehrseffizienz im innerstädtischen Raum signifikant zu steigern. So können Verdeckungseffekte, wie sie beispielsweise durch Gebäude oder andere Verkehrsteilnehmer entstehen, aufgelöst werden. Darüber hinaus wird eine kooperative Manöverplanung ermöglicht, bei der die Bewegung mehrerer vernetzter automatisierter Fahrzeuge gemeinsam optimiert wird. Viele existierende Ansätze für das automatisierte Kreuzungsmanagement betrachten lediglich den vollständig automatisierten Verkehr. In der Praxis wird Mischverkehr, also die Interaktion von menschlichen Fahrern und automatisierten Fahrzeugen, vorherrschend sein. Im vorliegenden Beitrag wird ein Verfahren basierend auf Reinforcement Learning und Graph Neural Networks für das automatisierte Kreuzungsmanagement im Mischverkehr vorgestellt. In simulativen Experimenten zeigt der Ansatz gegenüber einer regelbasierten Baseline einen deutlichen Gewinn an Effizienz. Zudem werden Experimente mit einem realen Versuchsträger auf einer Erprobungsbahn vorgestellt.

Schlüsselwörter: Kooperative Manöverplanung, Reinforcement Learning, Graph Neural Networks, vernetztes automatisiertes Fahren

1 Einleitung

Der innerstädtische Verkehr leidet vielerorts unter Störungen und Ineffizienzen aufgrund von stetig steigendem Fahrzeugaufkommen, dem klassische Verkehrsmanagementsysteme nicht gewachsen sind. Insbesondere kleinere Kreuzungen werden in der Regel durch statische Vorrangregeln koordiniert, sodass von einer Nebenstraße kommende Fahrzeuge Vorrang gewähren müssen. Darüber hinaus verdecken Gebäude und andere Verkehrsteilnehmer häufig die Sicht auf den kreuzenden Verkehr sowohl für menschliche Fahrer als auch fahrzeuggebundene Sensorsysteme.

Der zunehmende Einsatz von vernetzten Fahrzeugen (VFs) und vernetzten automatisierten Fahrzeugen (VAFs) eröffnet neue Optionen, die Verkehrseffizienz zu verbessern. Diese Fahrzeuge sind über eine drahtlose Kommunikationsschnittstelle vernetzt und können neben der Ankündigung ihrer Präsenz auch Perzeptionsdaten teilen. Durch die zunehmende Verfügbarkeit von Edge-Computing-Ressourcen kann darüber hinaus ein lokales Umfeldmodell, beispielsweise einer Kreuzung und deren Umgebung, an zentraler Stelle vorgehalten werden. Ein solches Umfeldmodell kann durch einen Edge-Server an VFs im

^{*}Robert Bosch GmbH, Corporate Research, D-71272 Renningen (E-Mail: {marvin.klimke, benjamin.volz}@de.bosch.com).

[†]Institut für Mess-, Regel- und Mikrotechnik – Universität Ulm, Albert-Einstein-Allee 41, D-89081 Ulm (E-Mail: michael.buchholz@uni-ulm.de).

Betriebsbereich verteilt werden. Diese zusätzlichen Informationen können zur Verbesserung der lokalen Planung auf einem VAF verwendet werden.

In dem öffentlich geförderten Projekt MEC-View wurde das Potential des vernetzten automatisierten Fahrens an einer als Pilotanlage ausgestatteten suburbanen Dreiwegekreuzung in Ulm-Lehr erforscht [1]. Die Sicht auf die vorfahrtsberechtigten Straße ist für Fahrzeuge auf der Nebenstraße durch Bebauung stark eingeschränkt. Nur mit Unterstützung durch das serverseitige Umfeldmodell ist es einem automatisierten Fahrzeug möglich, sich ohne starke Verzögerung in eine existierende Lücke im priorisierten Verkehr einzufügen.

Die vorliegende Arbeit baut auf diesem Ansatz auf und untersucht das Potential einer zentralisierten, kooperativen Manöverplanung durch den Einsatz von Reinforcement Learning (RL). Auf Basis des serverseitigen Umfeldmodells wird ein gemeinsamer Handlungsplan für alle VAFs in der Verkehrsszene abgeleitet, der den Fahrzeugen als Handlungsempfehlung übermittelt wird. Durch die Einbeziehung der VAFs auf der Hauptstraße werden explizite Abweichungen von den Vorrangregeln möglich. Das Einfädeln eines VAFs von der Nebenstraße kann beispielsweise durch proaktives Abbremsen eines VAFs auf der Hauptstraße ermöglicht werden. Auf öffentlichen Straßen wird die Durchdringung mit VFs und VAFs in absehbarer Zeit nicht annähernd 100 % erreichen. Stattdessen werden menschliche Fahrer und automatisierte Fahrzeuge im Mischverkehr interagieren. Daher müssen in der kooperativen Manöverplanung alle Fahrzeuge berücksichtigt werden, einschließlich solcher, die nicht beeinflusst werden können.

Im vorliegenden Beitrag wird in Kapitel 2 zunächst das RL-Modell und das Graph Neural Network (GNN) für die kooperative Manöverplanung im Mischverkehr aus [2] vorgestellt. Anschließend werden Ergebnisse der simulativen Evaluation im Vergleich mit einer regelbasierten Baseline präsentiert (Kap. 3). Kapitel 4 behandelt die Anbindung des Lernmodells an einen Bewegungsplanungsalgorithmus und zeigt Experimente unter Verwendung eines realen Versuchsträgers. Die vorliegende Arbeit wird in Kapitel 5 zusammengefasst.

2 Reinforcement Learning Modell für Kooperative Manöverplanung

Dieses Kapitel behandelt das vorgeschlagene Lernmodell (Abschnitt 2.1), die Graph-basierende Eingaberepräsentation (Abschnitt 2.2), die Netzarchitektur (Abschnitt 2.3) sowie die Gestaltung der Reward-Funktion (Abschnitt 2.4).

2.1 Lernmodell

Die kooperative Planung über mehrere Fahrzeuge wird als Mehragenten-RL-Problem aufgefasst. Aufgrund der zentralisierten Planung kann das Problem als einzelner, teilweise beobachtbarer Markov-Entscheidungsprozess (POMDP) formuliert werden:

$$(S, A, T, R, \Omega, O). \quad (1)$$

Die Menge der Zustände S beinhaltet alle erreichbaren Simulatorzustände einschließlich vollständiger Zustandsinformation aller automatisierter und manuell gefahrener Fahrzeuge. Während der Aktionsraum durch die Menge A definiert wird, ist die bedingte

Wahrscheinlichkeit eines Zustandswechsels von $s \in S$ nach $s' \in S$ unter Anwendung von Aktion $a \in A$ durch $T(s' | s, a)$ gegeben. Da die Reaktion menschlicher Fahrer nichtdeterministisch ist, sind auch die Zustandsübergänge nichtdeterministisch. $R : S \times A \rightarrow \mathbb{R}$ beschreibt die Reward-Funktion, die eine gewählte Aktion in einem gegebenen Zustand in Form eines Skalars bewertet. Da der Großteil des abstrakten Verkehrszustandes S für den kooperativen Planer nicht beobachtbar ist, wird eine reduzierte Beobachtungsmenge Ω definiert. Die Wahrscheinlichkeit, dass ein Zustand $s \in S$ in der Beobachtung $\omega \in \Omega$ resultiert, wird durch $O(\omega | s)$ angegeben. Die Abbildung ist nicht injektiv, was sich beispielsweise dann äußert, wenn der Fahrer eines nicht vernetzten Fahrzeugs seine Manöverintention ohne extern erkennbaren Hinweis ändert. Die Dimensionalität der Zustands- und Aktionsräume hängt von der aktuellen Anzahl (steuerbarer) Fahrzeuge ab und kann im Laufe der Zeit variieren.

2.2 Graph-basierende Eingaberepräsentation

Eine geeignete Eingaberepräsentation ist elementar für den erfolgreichen Einsatz von künstlichen neuronalen Netzen. Für die Eingaberepräsentation bei der kooperativen Manöverplanung sind dabei folgende drei zentralen Anforderungen zu berücksichtigen:

- Invarianz bezüglich der Anzahl der Fahrzeuge in der Szene,
- Permutationsinvarianz hinsichtlich der Eingabeknoten,
- Permutationsäquivarianz bezüglich der Ein- und Ausgabeknoten.

Einfache tabellarische Repräsentationen verletzen bereits die Invarianzeigenschaften. Die Beschränkungen von Eingängen mit fester Größe werden in [3] genauer diskutiert. Ein gerendertes Rasterbild der Verkehrsszene, wie es häufig mit Faltungsnetzen verwendet wird, erfüllt die Invarianzeigenschaften, aber erfordert typischerweise ein Ego-Fahrzeug, auf das es zentriert wird [4]. Dieser Prozess müsste für jedes beteiligte Fahrzeug wiederholt werden, was eine Anwendung auf komplexe Verkehrsszenen rechnerisch kaum durchführbar macht. Permutationsäquivarianz beschreibt die Tatsache, dass die abgeleiteten Ausgaben für jedes Fahrzeug unabhängig von der Reihenfolge in der Eingaberepräsentation sind. In der vorliegenden Arbeit wird daher eine flexible Graph-basierende Repräsentation vorgeschlagen.

Folglich wird eine Beobachtung beschrieben durch $\omega = (V, E, U)$, wobei jedes Fahrzeug auf einen Knoten $\nu \in V$ abgebildet wird, die Menge der Kanten durch E und die Menge der Kantentypen durch U bezeichnet wird. Die verfügbaren Kantentypen ergeben sich aus dem kartesischen Produkt aus den Dimensionen Konfliktrelation und Automatisierungstyp der beteiligten Fahrzeuge, wie in Abbildung 1 dargestellt. Formal ist eine Kante definiert als $(\nu_i, \nu_j, g_{ij}, r) \in E$, wobei die Quell- und Zielknoten als ν_i respektive ν_j und der Kantentyp als r bezeichnet werden. Nach [2] werden die Position, Geschwindigkeit, gemessene Beschleunigung sowie ein Indikator, ob es sich um ein VAF handelt, als eingeangssseitige Knoteneigenschaften $\mathbf{h}^{(0)}$ encodiert. Die Kanteneigenschaften $\mathbf{g}_{ij}^{(0)}$ enthalten eine Distanzmetrik sowie eine relative Peilung und die Vorrangrelation zweier Fahrzeuge.

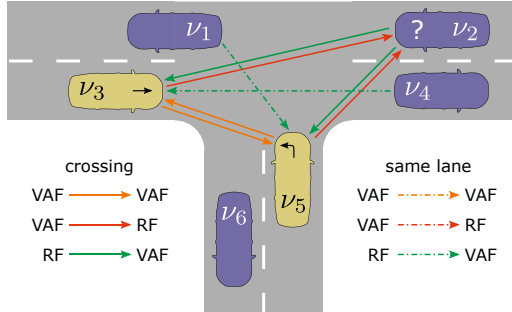


Abbildung 1: Die Graph-basierende Eingaberepräsentation gemäß [2]. Die Route eines VAF (gelb) ist durch einen Pfeil gekennzeichnet. Aufgrund der unbekannten Route des regulären Fahrzeugs ν_2 (RF, blau, gekennzeichnet mit '?') teilt es Kanten mit beiden VAFs, obwohl der Konflikt mit ν_3 nur bei einem Abbiegen von ν_2 auftreten wird.

2.3 Architektur des Graph Neural Networks

Für das Training des GNN wird der RL-Algorithmus TD3 aus [5] verwendet, der zur Familie der Actor-Critic-Methoden für kontinuierliche Aktionsräume gehört. Da sich das Actor- und das Critic-Netz nur unwesentlich auf der Ausgabeseite unterscheiden, wird im Folgenden lediglich die Actor-Netzarchitektur im Detail vorgestellt.

In diesem Beitrag wird eine GNN-Architektur, bestehend aus Schichten vom Typ Relational Graph Convolutional Network (RGCN) [6] sowie Graph Attention Network (GAT) [7], vorgeschlagen. Um neben Knoteneigenschaften auch Kanteneigenschaften verarbeiten zu können, wird in den RGCN-Schichten eine modifizierte Update-Vorschrift

$$\mathbf{h}_i^{(2)} = \sigma \left(\sum_{r \in U} \max_{j \in \mathcal{N}_i^r} \mathbf{W}_r^{(1)} [\mathbf{h}_j^{(1)}, \mathbf{g}_{ji}^{(1)}] + \mathbf{W}_0^{(1)} \mathbf{h}_i^{(1)} \right) \quad (2)$$

verwendet. Die unterschiedlichen Kantentypen, wie in Abschnitt 2.2 eingeführt, korrespondieren zu individuell gelernten Gewichtsmatrizen $\mathbf{W}_r$ je Kantentyp r . Dies ermöglicht dem GNN die fundamentalen Unterschiede in der Interaktion automatisierter und nicht-automatisierter Fahrzeuge auf sinnvolle Aktionen zu übertragen.

Abbildung 2 stellt die Netzarchitektur schematisch dar. Sowohl die Knoten- als auch die Kanteneigenschaften werden zunächst durch vollständig verbundene Schichten encodiert. Anschließend werden die latenten Knoteneigenschaften durch drei GNN-Schichten für den sogenannten Nachrichtenaustausch geleitet, welche die Kanteneigenschaften jeweils als zusätzliche Eingabe erhalten. Alle GNN-Schichten erfüllen die Vorschrift $\text{conv} : V^n \times E^m \rightarrow V^n$, wobei n die Anzahl der Fahrzeuge (Knoten) und m die Anzahl der Konflikte (Kanten) bezeichnet. Die Kanteneigenschaften werden durch das Netz nicht aktualisiert, da auf den Kanten keine Ausgabe zu inferieren ist. Die Benutzung alternierender Schichttypen zeigte bessere Ergebnisse als ein reines RGCN- oder GAT-Netz. Damit wird der Attention-Mechanismus der GAT-Schicht mit expliziten Kanteneigenschaften und Kantentypen verbunden. Insbesondere nutzen die GAT-Schichten selbst die Kantentypen nicht und verwenden die Kanteneigenschaften lediglich für die Attention-Gewichte,

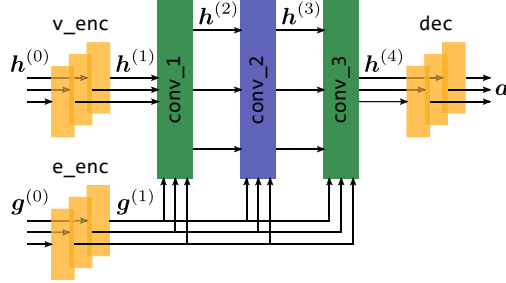


Abbildung 2: Die GNN-Architektur des Actor-Netzes aus [2]. Knoteneigenschaften $\mathbf{h}^{(0)}$ und Kanteneigenschaften $\mathbf{g}^{(0)}$ werden auf eine gemeinsame Aktion $\mathbf{a}$ abgebildet. Die um Unterstützung für Kanteneigenschaften erweiterten RGCN-Schichten sind in grün abgebildet, die GAT-Schicht in blau und die vollständig verbundenen Schichten in gelb.

aber nicht für das Update der Knotenfeatures. Die latenten Knoteneigenschaften $\mathbf{h}^{(4)}$ werden dann durch weitere vollständig verbundene Schichten geleitet, um eine gemeinsame Aktion für alle Fahrzeuge in der Szene abzuleiten. Alle verdeckten Schichten werden mit einer ReLU-Aktivierungsfunktion versehen und alle vollständig verbundenen Schichten teilen ihre Gewichte über Knoten beziehungsweise Kanten. Im Gegensatz zum Actor-Netz werden im Critic-Netz die latenten Knoteneigenschaften $\mathbf{h}^{(4)}$ zu einem einzelnen Vektor aggregiert, aus dem sodann eine Schätzung des Q -Wertes abgeleitet wird.

2.4 Reward-Funktion

Die Reward-Funktion für die Steuerung des RL-Trainings ist als gewichtete Summe

$$R = \sum_{k \in \mathcal{R}} w_k R_k \quad (3)$$

definiert, wobei die Gewichte mit w_k und die Menge der Reward-Komponenten mit $\mathcal{R}$ bezeichnet werden.

Die treibende positive Reward-Komponente begünstigt hohe mittlere Fahrgeschwindigkeiten mit einem linearen Anstieg bis zur gültigen Geschwindigkeitsbegrenzung. Ein triviales kollisionsfreies Verhalten wird erzielt, indem alle Fahrzeuge angehalten werden. Da dies keine zulässige Lösung für die kooperative Planung ist, wird dieser Fall durch einen Stillstand-Abschlag im Training vermieden. Um dem Modell das Einhalten sinnvoller Sicherheitsabstände beizubringen, wird eine geschwindigkeitsabhängige Distanzmetrik als weiterer Abschlag eingeführt. Simulationen im Mischverkehr zeigten ein prinzipiell sehr konservatives Verhalten der VAF, die potentiell jedes RF vorfahren lassen, selbst wenn das VAF Vorrang hat. Um dieses Verhalten zu korrigieren, wird eine weitere Reward-Komponente eingeführt, welche einen Abschlag auslöst, wenn ein VAF ohne erkennbaren Grund weit entfernt vom Kreuzungseingang zum Stehen kommt. Abschließend wird jede Kollision mit einem negativen Reward belegt, damit das Modell implizit Kollisionsvermeidung erlernt. Darüber hinaus wird die Trainingsepisode im Kollisionsfall abgebrochen.

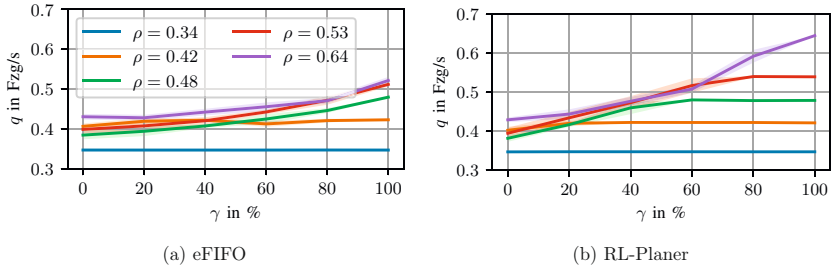


Abbildung 3: Erreichter Fahrzeugdurchsatz über variierenden Anteil an VAFs bei Einsatz der eFIFO-Vergleichsbasis und des RL-Planers. Jede Linienfarbe korrespondiert zu einem Szenario mit zunehmender Verkehrsnachfrage ρ in Fzg/s. Die schattierte Fläche markiert die Standardabweichung.

3 Simulative Evaluation

Das Training und die simulative Auswertung des RL-Planers erfolgen in der Open-Source-Simulationsumgebung Highway-env [8], die geringfügig für die Anwendung in kooperativer Manöverplanung angepasst wurde. Das Verhalten von Menschen gefahrener Fahrzeuge wird durch ein Fahrermodell, wie das Intelligent Driver Model (IDM) [9], simuliert. In der Evaluation wird abweichend ein erweitertes IDM [10] verwendet, welches nicht-deterministisches Verhalten abbildet. Darüber hinaus wurde die Vorranglogik in Highway-env verbessert, um auch bei hohen Verkehrsdichten einsetzbar zu sein.

In diesem Beitrag wird ein erweitertes First-In-First-Out-Schema (eFIFO) als Vergleichsbasis für die kooperative Manöverplanung verwendet. VAFs werden auf Basis ihrer Distanz zur Kreuzung priorisiert und Vorrang von RFs wird als Randbedingung berücksichtigt. Darüber hinaus können auch konfliktfreie Pfade zeitgleich befahren werden, was einen erheblichen Performance-Gewinn liefert.

Die kooperativen Planungsansätze wurden in fünf Szenarien mit unterschiedlicher Verkehrsnachfrage ρ an einer Vierwegekreuzung evaluiert. Dabei wurde jedes Szenario mit variierendem Anteil an VAFs im Gesamtverkehr γ mehrfach simuliert. Der tatsächlich erreichte Durchsatz wird durch die *Flow Rate* q erfasst. Wie in Abbildung 3 ersichtlich ist, erreicht der RL-Planer gegenüber der Vergleichsbasis einen deutlich höheren Fahrzeugdurchsatz, auch für geringe Ausstattungsdaten.

Die Generalisierung des Ansatzes auf andere Kreuzungen, die nicht im Training verwendet wurden, findet sich für vollständig automatisierten Verkehr in [11].

4 Experimente mit Bewegungsplanung

Das zentralisierte Planungsmodul im vernetzten automatisierten Fahren erstellt ein Verhaltensplan für alle beteiligten Fahrzeuge, sendet jedoch keine direkten Steuerkommandos an die VAFs. Stattdessen führt jedes VAF lokal einen Bewegungsplanungsalgorithmus aus, welcher die kooperativen Manöver in eine fahrbare Trajektorie umsetzt. Dafür ist ein

Vorausplanungshorizont von zumindest mehreren Sekunden erforderlich. Das RL-Modell liefert jedoch nur eine instantane Aktion für den nächsten Zeitschritt, bevor eine neue Zustandsbeobachtung eintrifft.

In der vorliegenden Arbeit wird der Einsatz einer serverseitig integrierten Simulationsumgebung nach [12] vorgeschlagen, um die Lücke zwischen iterativer RL-Planung und Vorausplanung einer Trajektorie zu schließen. Um ein kooperatives Manöver zu planen, wird der aktuelle Zustand des serverseitigen Umfeldmodells $\mathcal{E}$ als Startzustand in der Simulationsumgebung abgebildet, wie in Abbildung 4 dargestellt. In der Folge wird die Entwicklung der Verkehrsszene in die Zukunft prädiziert, wobei das RL-Modell in jedem Zeitschritt abgefragt wird, um das Verhalten der VAFs zu bestimmen. Die regulären Fahrzeuge in der Szene werden unter Annahme eines Fahrermodells, wie dem IDM, mit-prädiziert. Sobald alle Fahrzeuge in der Simulation ihre jeweilige Zielseite der Kreuzung erreicht haben, endet der Simulationslauf. Sodann wird aus der simulierten Trajektorie je Fahrzeug eine abstrakte Manöverdarstellung $\mathcal{O}$ abgeleitet, die den VAFs übermittelt wird.

Die in diesem Beitrag verwendete Darstellung basiert auf dem Vorschlag zur *Maneuver Coordination Message* aus [13]. Während die laterale Führung anhand einer Spurmittellinie erfolgt, werden die Fahrzeuge longitudinal durch zusätzliche Manöverwegpunkte beeinflusst. Neben einer Position ist ein Manöverwegpunkt durch Einschränkungen des Zeitpunktes δt_{MWP} und der Geschwindigkeit bei der Überquerung gekennzeichnet. Jedes VAF erhält in einem kooperativen Manöver einen Manöverwegpunkt am Kreuzungseingang. Durch Anpassung der Zeit- beziehungsweise Geschwindigkeitsvorgaben wird die Überquerung der Kreuzung durch alle beteiligten VAF koordiniert.

In Vorbereitung auf die Fahrttests mit mehreren VAFs im Realverkehr wurde der vorgeschlagene Ansatz zunächst in einem automatisierten Fahrzeug auf einer Erprobungsbahn getestet. Das Experiment folgt dem Prinzip einer *Vehicle-in-the-Loop*-Simulation, bei der das automatisierte Fahrzeug in der Realität bewegt wird und zusätzliche Objektfahrzeuge simuliert werden können. Im untersuchten Manöver quert das automatisierte Fahrzeug eine Kreuzung, an der es einem Objektfahrzeug Vorrang gewähren muss. Dabei werden zwei alternative Überplanungskonzepte evaluiert. Zunächst wird in einem idealisierten Modus das Manöver *einmalig* geplant und fortan unverändert ausgeführt. Insbesondere im Mischverkehr, unter Anwesenheit nicht beeinflussbarer Fahrzeuge, kann sich die tatsächliche Entwicklung des Manövers von dem ursprünglichen Plan unterscheiden. Um diese Fälle abzufangen, wird eine *zyklische* Neuplanung vorgeschlagen, wobei das Manöver in diesen Experimenten mit

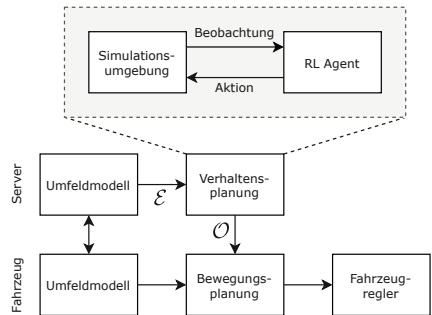


Abbildung 4: Verwendung einer serverseitig integrierten Simulationsumgebung für die kooperative Manöverplanung nach [12]. Aus dem serverseitigen Umfeldmodell $\mathcal{E}$ wird eine abstrakte Manöverdarstellung $\mathcal{O}$ für jedes VAF abgeleitet.

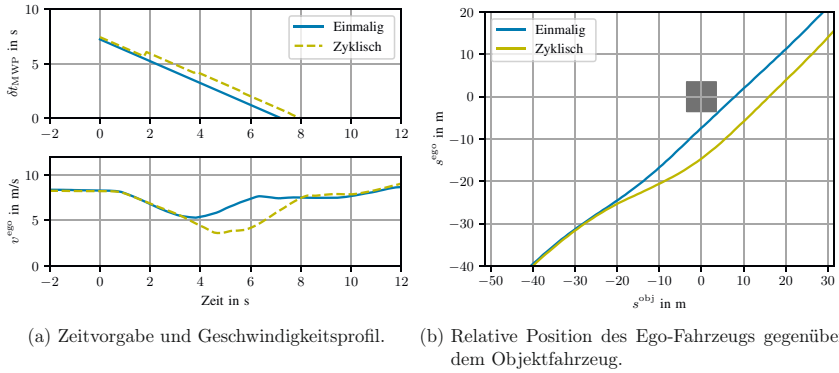


Abbildung 5: Exemplarische Aufzeichnung eines Manövers bei einmaliger und zyklischer Planung aus [12]. Bei zyklischer Neuplanung wird die Überquerung der Kreuzung mehrfach verzögert. Das graue Rechteck stellt den Bereich dar, in dem es zur Kollision kommt.

einer Periodenlänge von 2 s neu geplant wird. Dieser Wert bezieht sich lediglich auf die serverseitige RL-basierte Verhaltensplanung, während die unterlagerte Bewegungsplanung im Fahrzeug die Trajektorie wesentlich hochfrequenter plant.

Eine exemplarische Aufzeichnung eines solchen Manövers ist in Abbildung 5 visualisiert. In diesem Fall fährt das Objektfahrzeug langsamer als die in der Karte hinterlegte Maximalgeschwindigkeit, die in der Verhaltensplanung als Zielgeschwindigkeit angenommen wird. Bei zyklischer Neuplanung wird die Zeitvorgabe δt_{MWP} des Manöverwegpunkts daher mehrfach verschoben, was sich durch eine stärkere Verzögerung auch im Geschwindigkeitsprofil bemerkbar macht. Bei Betrachtung von Abbildung 5(b) wird deutlich, dass diese Verzögerung erforderlich ist, um einen hinreichenden Sicherheitsabstand zum Objektfahrzeug zu wahren. Unter der Annahme von rechteckigen Fahrzeugausmaßen von $5 \text{ m} \times 2 \text{ m}$ kollidieren die Fahrzeuge, wenn die Trajektorie den grauen Bereich durchschreitet. In diesem Beispiel verbleibt bei einmaliger Planung keinerlei Abstand zwischen den Fahrzeugen, was in der Praxis nicht akzeptabel ist. Nur durch zyklische Neuplanung wird die geringere Geschwindigkeit des Objektfahrzeugs korrekt antizipiert und ein hinreichender Sicherheitsabstand eingehalten.

Für eine genauere Untersuchung dieses Effekts wurden 40 komplexere Szenarien mit bis zu sechs VAFs simulativ mit einmaliger und zyklischer Planung evaluiert. Dabei wird

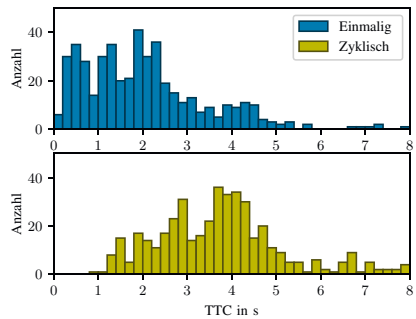


Abbildung 6: Verteilung der TTC über 40 simulierte Szenarien in einmaliger und zyklischer Planung.

eine verallgemeinerte Time-to-Collision (TTC) verwendet, die auch für kreuzende Pfade anwendbar ist und in [14] vorgeschlagen wurde. Abbildung 6 zeigt die Verteilung der TTC für beide Planungsmodi als Histogramm. Offensichtlich treten bei einmaliger Planung wesentlich häufiger kritische TTC von unter 1 s auf. In der Simulation mit einmaliger Planung endeten sechs Szenarien in einer Kollision, die durch eine TTC von 0 s im Plot wiederzufinden sind. Diese simulativen Experimente wurden im vollständig automatisierten Verkehr durchgeführt und es ist zu erwarten, dass sich der Effekt im Mischverkehr noch stärker ausgeprägt. Auf Basis dieser initialen Experimente auf der Erprobungsbahn und im Simulator erscheint der zyklische Planungsmodus prinzipiell geeignet, um die RL-basierte Verhaltensplanung an die fahrzeugseitige Bewegungsplanung anzubinden.

5 Zusammenfassung

In diesem Beitrag wurde ein auf maschinellem Lernen basierender Ansatz für die Planung kooperativer Manöver im urbanen Mischverkehr vorgestellt. Durch Verwendung einer Graph-basierten Eingaberepräsentation können komplexe Verkehrsszenen an Kreuzungen effizient encodiert werden. Aufgrund fehlender Trainingsdaten für kooperative Manöver wird das GNN durch Reinforcement Learning auf einer geeigneten Reward-Funktion trainiert. Auswertungen in der Simulationsumgebung Highway-env zeigen eine deutliche Verbesserung des Verkehrsdurchsatzes auch bei geringen Ausstattungsraten im Mischverkehr. Auf einer Erprobungsbahn wurden erste Experimente in einem Versuchsträger mit dedizierter Bewegungsplanung durchgeführt und gezeigt, dass unter zyklischer Neuplanung ausgewählte Kreuzungsszenarien beherrscht werden.

Danksagung

Teile dieser Arbeit wurden im Rahmen des Projekts LUKAS (Förderkennzeichen 19A2000 4A und 19A20004F) durchgeführt, welches durch das Programm “Hoch- und vollautomatisiertes Fahren in anspruchsvollen Fahrsituationen” des Bundesministeriums für Wirtschaft und Klimaschutz finanziert wird.

Literatur

- [1] M. Buchholz, J. Müller, M. Herrmann, J. Strohecker, B. Völz, M. Maier, J. Paczia, O. Stein, H. Rehborn, and R.-W. Henn, “Handling Occlusions in Automated Driving Using a Multiaccess Edge Computing Server-Based Environment Model From Infrastructure Sensors,” *IEEE Intelligent Transportation Systems Magazine*, vol. 14, no. 3, pp. 106–120, 2022.
- [2] M. Klimke, B. Völz, and M. Buchholz, “Automatic Intersection Management in Mixed Traffic Using Reinforcement Learning and Graph Neural Networks,” in *2023 IEEE Intelligent Vehicles Symposium (IV)*, to appear.
- [3] M. Huegle, G. Kalweit, B. Mirchevska, M. Werling, and J. Boedecker, “Dynamic Input for Deep Reinforcement Learning in Autonomous Driving,” in *2019 IEEE/RSJ*

- International Conference on Intelligent Robots and Systems (IROS)*, pp. 7566–7573, 2019.
- [4] J. Chen, B. Yuan, and M. Tomizuka, “Deep Imitation Learning for Autonomous Driving in Generic Urban Scenarios with Enhanced Safety,” in *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pp. 2884–2890, 2019.
 - [5] S. Fujimoto, H. van Hoof, and D. Meger, “Addressing Function Approximation Error in Actor-Critic Methods,” in *Proceedings of the 35th International Conference on Machine Learning* (J. Dy and A. Krause, eds.), vol. 80 of *Proceedings of Machine Learning Research*, pp. 1587–1596, 2018.
 - [6] M. Schlichtkrull, T. N. Kipf, P. Bloem, R. van den Berg, I. Titov, and M. Welling, “Modeling Relational Data with Graph Convolutional Networks,” in *The Semantic Web* (A. Gangemi, R. Navigli, M.-E. Vidal, P. Hitzler, R. Troncy, L. Hollink, A. Tor-dai, and M. Alam, eds.), vol. 10843, pp. 593–607, Springer International Publishing, 2018.
 - [7] P. Velickovic, G. Cucurull, A. Casanova, A. Romero, P. Lio, and Y. Bengio, “Graph Attention Networks,” in *International Conference on Learning Representations*, 2018.
 - [8] E. Leurent, “An Environment for Autonomous Driving Decision-Making,” 2018.
 - [9] M. Treiber, A. Hennecke, and D. Helbing, “Congested Traffic States in Empirical Observations and Microscopic Simulations,” *Physical Review E*, vol. 62, no. 2, pp. 1805–1824, 2000.
 - [10] D. Salles, S. Kaufmann, and H.-C. Reuss, “Extending the Intelligent Driver Model in SUMO and Verifying the Drive Off Trajectories with Aerial Measurements,” in *SUMO User Conference*, 2020.
 - [11] M. Klimke, J. Gerigk, B. Völz, and M. Buchholz, “An Enhanced Graph Representation for Machine Learning Based Automatic Intersection Management,” in *2022 IEEE International Conference on Intelligent Transportation Systems (ITSC)*, pp. 523–530, 2022.
 - [12] M. Klimke, B. Völz, and M. Buchholz, “Integration of Reinforcement Learning Based Behavior Planning With Sampling Based Motion Planning for Automated Driving,” in *2023 IEEE Intelligent Vehicles Symposium (IV)*, to appear.
 - [13] M. B. Mertens, J. Müller, R. Dehler, M. Klimke, M. Maier, S. Gherekhloo, B. Völz, R.-W. Henn, and M. Buchholz, “An Extended Maneuver Coordination Protocol with Support for Urban Scenarios and Mixed Traffic,” in *2021 IEEE Vehicular Networking Conference (VNC)*, pp. 32–35, 2021.
 - [14] F. Jimenez, J. E. Naranjo, and F. Garcia, “An Improved Method to Calculate the Time-to-Collision of Two Vehicles,” *International Journal of Intelligent Transportation Systems Research*, vol. 11, no. 1, pp. 34–42, 2013.

Semantische Normverhaltensanalyse zur durchgängigen, formalen Verhaltensspezifikation automatisierter Straßenfahrzeuge

Nayel Fabian Salem^{*†}, Veronica Haber[‡], Marcus Nolte[†],
Robert Graubohm[†], Matthias Rauschenbach[§], Jan Reich[¶],
Torben Stolte[†] und Markus Maurer[†]

Zusammenfassung: Die Absicherung automatisierter Straßenfahrzeuge (SAE Level 3+) setzt die Spezifikation und Überprüfung des Verhaltens eines Fahrzeugs in seiner Betriebsumgebung voraus. Mithilfe von Szenarien kann der offene Verkehrskontext, in dem ein solches System agiert, strukturiert beschrieben werden. Um Annahmen, welche bei der Verhaltensspezifikation innerhalb von Szenarien getroffen werden, begründen und belegen zu können, ist eine durchgängige Dokumentation von Entwurfsentscheidungen erforderlich. In dieser Arbeit wird die *semantische Normverhaltensanalyse* vorgestellt, mithilfe derer Ansprüche an das Verhalten eines automatisierten Fahrzeugs in seiner Betriebsumgebung durchgängig auf ein formales Regelsystem aus semantischen Konzepten für ausgewählte Szenarien abgebildet werden können. Der Beitrag der vorgestellten Methode besteht in der Rückverfolgbarkeit dieser formalisierten Konzepte zu den damit verbundenen Ansprüchen an das Verhalten. Die Durchführung einer semantischen Normverhaltensanalyse wird beispielhaft an zwei Szenarien demonstriert.

Schlüsselwörter: Durchgängigkeit, Verhaltensspezifikation, Wissensrepräsentation

1 Einleitung

Die fortschreitende Automatisierung der Fahraufgabe stellt Beteiligte aus Gesetzgebung, Technik und Ethik gleichermaßen vor neue Herausforderungen. Um die mit dem automatisierten Fahren (SAE Level 3+ [1]) Versprechungen – z.B. die Erhöhung der Verkehrssicherheit – zu erfüllen, müssen Anforderungen an sicheres und regelkonformes Verhalten automatisierter Straßenfahrzeuge im Straßenverkehr formuliert und geprüft werden.

Für Entwickler*innen automatisierter Straßenfahrzeuge ergibt sich daraus unter anderem die Aufgabe, zu begründen und zu kommunizieren, warum sich das Fahrzeug regelkonform und sicher im Straßenverkehr verhalten wird. Um regelkonformes Verhalten argumentieren und belegen zu können, ist es notwendig, Anforderungen an regelkonformes Verhalten zunächst zu formulieren und anschließend zu überprüfen. Für die Entwicklung

^{*}Korrespondierender Autor: salem@ifr.ing.tu-bs.de

[†]Institut für Regelungstechnik, Technische Universität Braunschweig, Braunschweig

[‡]PROSTEP AG, München

[§]Fraunhofer-Institut für Betriebsfestigkeit und Systemzuverlässigkeit LBF, Darmstadt

[¶]Fraunhofer-Institut für Experimentelles Software Engineering IESE, Kaiserslautern

automatisierter Straßenfahrzeuge bedeutet dies – insbesondere für eine Konformität mit der ISO 21448 [2] – eine Spezifikation des zu implementierenden Verhaltens in Bezug zu seiner Betriebsumgebung.

Eine besondere Herausforderung bei der Verhaltensspezifikation automatisierter Straßenfahrzeuge ergibt sich aus dem Betrieb dieser Systeme im offenen Verkehrskontext, der bedingt, dass szenarienbasierte Ansätze der Absicherung inhärent mit epistemischen und aleatorischen Unsicherheiten einhergehen [3]. Die im Folgenden vorgestellte *semantische Normverhaltensanalyse* unterstützt den Prozess der expliziten Dokumentation von Annahmen und getroffenen Entwurfsentscheidungen bei der Spezifikation von zu implementierendem Verhalten durch die Formalisierung von Verhaltensregeln und kann dadurch die Nachvollziehbarkeit residualer Unsicherheiten erhöhen. Die Vollständigkeit der resultierenden Verhaltensspezifikation ist in Bezug auf den offenen Kontext weiterhin inhärent nicht gegeben, da die Analyse sich auf einen (idealerweise repräsentativen) Szenarienkatalog bezieht. Formalisierte Ansätze zur Verhaltensbeschreibung ermöglichen jedoch, innerhalb der verwendeten Formalismen, Konsistenz zu belegen.

Diese Arbeit beschreibt zunächst die verwendete Terminologie zur Verhaltensspezifikation (Abschnitt 2) und verwandte Arbeiten (Abschnitt 3). Anschließend erläutern wir die Bestandteile der semantischen Normverhaltensanalyse (Abschnitt 4). Die Anwendung der Methode wird an einem Beispiel illustriert (Abschnitt 5) und evaluiert (Abschnitt 6).

2 Terminologie

Im Folgenden definieren wir das *Normverhalten* als das aus gesetzlichen, gesellschaftlichen und ethischen sowie sicherheitsbezogenen Ansprüchen resultierende Verhalten eines Akteurs in einem *szenarienübergreifenden Kontext*.

Dagegen ist das *Sollverhalten* das sich aus gesetzlichen, gesellschaftlichen und ethischen sowie sicherheitsbezogenen Ansprüchen abgeleitete, umzusetzende Verhalten eines Akteurs in einem *szenarienspezifischen Kontext*.

Der Begriff *Verhalten* folgt hierbei im Gegensatz zu Censi u. a. [4] der Definition im Rahmen der verhaltensbasierten Robotik [5, 6] und beinhaltet dementsprechend neben dem von außen beobachtbaren Verhalten auch die dafür notwendigen Informationsflüsse als Stimuli der Systemantwort des Akteurs.

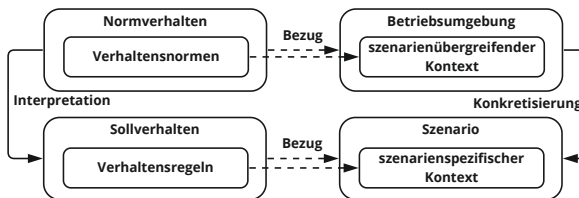


Abbildung 1: Darstellung der Beziehung von Normverhalten und Sollverhalten.

Im Folgenden werden allgemeine Verhaltensnormen in Summe als Normverhalten und konkrete Verhaltensregeln als Sollverhalten bezeichnet (Abbildung 1). Die Berücksichtigung gegebenenfalls konfliktärer Verhaltensnormen in einem szenarienübergreifenden

Kontext beziehungsweise einer Betriebsumgebung [7] (Normverhalten) erfordert bei der Formulierung von Verhaltensregeln in einem gegebenen Szenario [8] (Sollverhalten) sowohl eine szenarienspezifische Interpretation als auch eine Auflösung und damit verbundene Abwägung von Zielkonflikten.

Als Verhaltensnorm könnte beispielsweise formuliert werden, dass durchgezogene Fahrstreifenmarkierungen im Allgemeinen nicht zu überfahren sind. In einem Szenario, in dem der eigene Fahrstreifen belegt und der benachbarte Fahrstreifen frei ist, könnte dagegen die Entwurfsentscheidung getroffen werden, eine Verhaltensregel zu formulieren, welche das Überfahren der durchgezogenen Markierung (z. B. aufgrund eines Notstands nach § 16 OWiG) begründet. Die Legitimität der Begründung wird gegebenenfalls *ex post* durch entsprechende Institutionen wie Gerichte und Ethikräte festzustellen sein. Eine Repräsentation dieser mit der Verhaltensspezifikation verbundenen Annahmen und Abwägungen ist für Entwickler*innen insbesondere notwendig, um eine Argumentationsgrundlage gegenüber firmeninternen und -externen Institutionen zu besitzen, falls nicht regelkonformes Verhalten eines automatisierten Fahrzeugs im Feld beobachtet wird.

Unter Verwendung der vorgestellten Terminologie präsentieren wir folgende Beiträge:

- Analysen von Verhaltensnormen eines szenarienübergreifenden Kontextes und dabei getroffene Annahmen werden explizit dokumentiert.
- Verhaltensnormen werden über eine szenarienbasierte Betrachtung durchgängig als formale Verhaltensregeln abgebildet.
- Der spezifizierte Regelkatalog kann hinsichtlich seiner Konsistenz durch die Anwendung formaler Logik überprüft werden.
- Die durchgängige Spezifikation von Verhalten ermöglicht eine Überprüfung der Konformität dokumentierter Verhaltensregeln mit analysierten Verhaltensnormen innerhalb betrachteter Szenarien.

3 Verwandte Arbeiten

Bezogen auf normative Randbedingungen ist die semantische Normverhaltensanalyse wie folgt einzuordnen: In der ISO 21448 wird die Erstellung einer *funktionalen Spezifikation* (engl. *functional specification*) gefordert, welche unter anderem die *Fahrregeln* (engl. *driving policy*) beinhaltet. Die semantische Normverhaltensanalyse ermöglicht eine explizite Dokumentation von Entwurfsentscheidungen bei der Entwicklung solcher Fahrregeln (bzw. des Sollverhaltens). Auch Analysen der funktionalen Sicherheit nach ISO 26262 wie eine Gefährdungsidentifikation können durch eine Spezifikation des Sollverhaltens unterstützt werden [9].

Zur Repräsentation spezifizierter Verhaltensregeln verwenden wir Ontologien und Regeln der Prädikatenlogik erster Ordnung [10]. Die Nutzung (semi-)formaler Sprachen erfüllt den Zweck der automatisierten Überprüfung der Verhaltensspezifikation hinsichtlich ihrer Konsistenz wie nach [11] und [12] gefordert. Diese automatisierbare Überprüfung ist eine Erweiterung gegenüber [13], wo mit einer strukturierten natürlichen Sprache bereits eine erste Formalisierung ohne einen Formalismus zur automatisierten Überprüfung erzeugt wird. Weitere Ansätze zur Formalisierung des Sollverhaltens [14–16] beschränken sich in der Verhaltensspezifikation auf die statische Umgebung.

In der wissensbasierten Szenariengenerierung [17, 18] ist die Auswahl relevanter Szenarien für die Absicherung automatisierter Fahrsysteme eine offene Herausforderung. Die semantische Normverhaltensanalyse kann zu dieser Auswahl beitragen, da sie die systematische Herleitung von Szenarien, in denen sich das Ego-Fahrzeug konform zur Verhaltensspezifikation verhält, unterstützt.

Im Fokus der hier vorgestellten Arbeit steht die Verhaltensspezifikation auf Ebene funktionaler Szenarien. Um Sollverhalten innerhalb von Szenarien dieser Abstraktionsebene beschreiben und analysieren zu können, wurde das Phänomen-Signal-Modell vorgestellt [19, 20]. Der Ansatz greift auf einen Satz formaler Regeln zurück, um Handlungsoptionen im Rahmen des Sollverhaltens zu schlussfolgern. Die semantischen Normverhaltensanalyse kann den Prozess zur Spezifikation dieser Verhaltensregeln unterstützen.

Dieser Anspruch unterscheidet sich von anderen Arbeiten, welche Ansätze zur Formalisierung, unter anderem von Verkehrsregeln, vorschlagen [21–25]. Die Motivation dieser Ansätze besteht in der direkten Übersetzung von Verkehrsregeln in formale Beschreibungssprachen. Dabei wird der Aspekt der Durchgängigkeit von Konzepten und Regeln in Bezug auf getroffene Entwurfsentscheidungen vernachlässigt. Im Fall von [21] und darauf aufbauenden Arbeiten [22–24] bleibt unklar, wie sich der Ansatz im Kontext einer szenarienbasierten Absicherung anwenden lässt. Der Ansatz nach [25] ordnet sich mit einem Fokus auf die Harmonisierbarkeit von Prozessen zur Formalisierung von Verkehrsregeln dem szenarienbasierten Ansatz zu, eine konkrete Umsetzung bleibt aber offen.

Insbesondere in der Luftfahrt [26] und der Robotik [27], aber auch im Rahmen des automatisierten Fahrens [4, 28], werden auch formale Modelle und Methoden zur Dokumentation des angestrebten Verhaltens automatisierter Systeme verwendet. Diese Ansätze unterscheiden sich insbesondere in ihrem expliziten Bezug (bzw. ihrer Durchgängigkeit) zu den zugrundeliegenden Quellen von Annahmen und Verhaltensanforderungen. Das automatisierte Fahren im offenen Verkehrskontext stellt mit der Moderation insbesondere gesellschaftlicher und rechtlicher Aspekte [4] aber auch im Umgang mit Unsicherheiten [3] Anforderungen an die Durchgängigkeit einer Verhaltensspezifikation, welche durch die untersuchte Literatur nur teilweise adressiert werden.

4 Ansatz zur semantischen Normverhaltensanalyse

Abbildung 2 stellt das Vorgehen der semantischen Normverhaltensanalyse dar. Im ersten Schritt werden die genutzten Wissensquellen für die Analyse des Normverhaltens gewählt (z.B. die StVO). Diese Wissensquellen können Informationen zu Konzepten und Regeln sowie Methoden zum Umgang mit den gewählten Wissensquellen beinhalten. Wissensquellen sind im Kontext der Betriebsumgebung auszuwählen und zu analysieren. Die Betriebsumgebung wird hier als Summe von Bedingungen verstanden, unter denen ein automatisiertes Fahrzeug erwartbar und zulässig betrieben wird [1].

Zur Modellierung des Norm- und Sollverhaltens müssen zunächst die Elemente, die für die formale Beschreibung der Verhaltensregeln notwendig sind, aus den Wissensquellen herausgearbeitet werden. Während in den meist natürlich-sprachlich dokumentierten Quellen auch implizites Wissen enthalten ist, kann mit Hilfe expliziter Konzepte erklärbar und rückverfolgbar geschlussfolgert werden. Zum Beispiel verwendet die StVO Konzepte wie *Fahrstreifen*. Dabei ist ein Fahrstreifen immer Bestandteil einer *Fahrbahn*. Das Wissen darüber, dass das Befahren eines Fahrstreifens einer Fahrbahn immer auch das Befahren

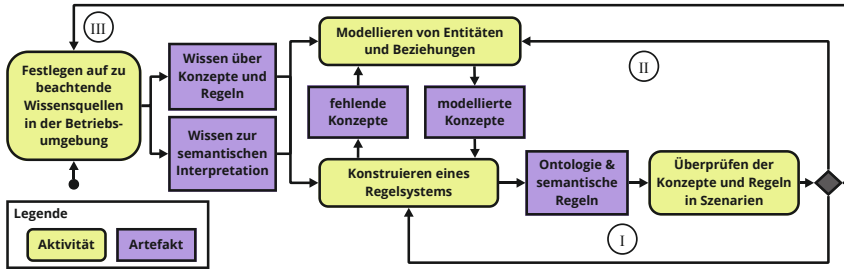


Abbildung 2: Darstellung des Vorgehens bei der semantischen Normverhaltensanalyse. Iterationen I, II und III werden durchlaufen, wenn die Überprüfung des inferierten Sollverhaltens Fehler im Regelsystem, Fehler bei den modellierten Konzepten oder Fehler bei der Auswahl der Wissensquellen feststellt. Fehlerhaftes Verhalten bedeutet: nonkonformes oder unsicheres Verhalten.

der Fahrbahn bedingt, erscheint zunächst trivial, kann aber bei der formalen Beschreibung eine wesentliche Schlussfolgerung sein.

Wissen über Konzepte und Regeln stammt häufig aus Wissensquellen, die Expertenwissen zur semantischen Interpretation voraussetzen. Daher sind zusätzlich Wissensquellen notwendig, die Methoden zum Umgang mit Domänenwissen beschreiben. Diese Methoden werden genutzt, um die Wissensquellen systematisch auf enthaltenes Wissen zu untersuchen und die Konzepte und Regeln zu formalisieren. Methodisches sowie domänenbezogenes Expertenwissen und damit verbundene semantische Interpretationen sind oft notwendig, weil Entitäten, Relationen und Regeln nicht immer explizit dokumentiert sind. Ohne Interpretation impliziter zu expliziten Konzepten wäre eine formale Repräsentation semantischer Beziehungen nicht möglich. Das in Abschnitt 5 gewählte Beispiel erfordert etwa die Repräsentation des Konzepts von „erkennbar querenden zu Fuß Gehenden“. Die StVO als Wissensquelle gibt keinen unmittelbaren Aufschluss darüber, welche Bedingungen an dieses Konzept geknüpft sind. Domänen-Expert*innen könnten hier zum Beispiel Gerichtsurteile heranziehen, um eine Interpretation zu stützen.

Eine Anbindung der Wissensmodellierung an existierendes Domänenwissen durch die flexible Wahl der Wissensquellen sowie der Methoden zur semantischen Interpretation ist ein zentrales Ziel der semantischen Normverhaltensanalyse. Der Ansatz stellt so die Anforderung an Domänen-Expert*innen, die Wissensquellen bezüglich enthaltener Konzepte und Regeln zu trennen. Die formale Strukturierung von Konzepten und Regeln ermöglicht eine formale Überprüfbarkeit der modellierten Verhaltensregeln hinsichtlich ihrer Konsistenz.

Zur formalen Repräsentation des natürlich-sprachlichen Wissens werden in der semantischen Normverhaltensanalyse Ontologien verwendet. Um diese Ontologien zur Verhaltensspezifikation nutzen zu können, werden Entitäten, deren Beziehungen sowie Inferenzregeln aus den zugrundeliegenden Wissensquellen modelliert. Die Konzeptualisierung von Entitäten und Beziehungen (Schritt 2) sowie die Konstruktion von Regeln (Schritt 3) bilden einen iterativen Prozess, da Konzepte in teilweise implizit repräsentiertem Wissen erst bei der Regelkonstruktion identifiziert werden können.

Verifiziert wird der erzeugte Regelkatalog in einem letzten Analyseschritt, in dem die Regeln auf alle betrachteten Szenarien angewendet werden. Dieser Schritt kann zum Beispiel mithilfe eines Phänomen-Signal-Modells [19, 20] durchgeführt werden. Mit der semantischen Normverhaltensanalyse wird zunächst nur eine systematische Ableitung von Verhaltensregeln für Szenen [8] eines funktionalen Szenarios [29, 30] beabsichtigt. Das Sollverhalten setzt sich auf der Beschreibungsebene eines funktionalen Szenarios aus den angewendeten Verhaltensregeln zusammen.

Aus der Verifikation ergeben sich im Rahmen der semantischen Normverhaltensanalyse (Abbildung 2) drei unterschiedliche Arten möglicher Fehlerfälle, die durch drei entsprechende Iterationsschleifen adressiert werden. Fehlerfälle bezeichnen eine Abweichung des aus dem spezifizierten Regelkatalog resultierenden Verhaltens von dem durch die Analyse der Wissensquellen vorgegebenen Verhalten und damit nonkonformes oder unsicheres Verhalten. Die erste Iteration (I) wird durch einen Fehlerfall aufgrund unzureichend definierter Regeln bedingt. Unzureichend bedeutet hier widersprüchlich oder nicht ausreichend detailliert. Ziel dieser Schleife ist die Erzeugung eines konsistenten Regelkatalogs. Der Eintritt in die zweite Iterationsschleife (II) wird durch fehlende Konzepte ausgelöst. Diese Schleife ist nötig, da selbst in einem konsistenten Regelkatalog festgestellt werden kann, dass die Konzeptualisierung für ein Szenario nicht hinreichend vollständig ist. Das Ziel der Schleife ist die Korrektheit des Regelkatalogs in Bezug auf die untersuchten Szenarien. In der dritten Schleife (III) kann eine in Bezug auf die Szenarien und Wissensquellen hinreichend vollständige Ontologie der Konzepte und ein widerspruchsfreier Regelkatalog vorliegen. Falls dennoch in Frage steht, ob das geschlussfolgerte, regelkonforme Verhalten dem tatsächlich gewünschten Sollverhalten entspricht, können die Wissensbasis geprüft und neue Quellen hinzugezogen werden.

Da Szenarien für die Ableitung von Verhaltensnormen genutzt werden, ist die Repräsentativität des Szenarienkatalogs maßgeblich verantwortlich für die Validität des formalisierten Sollverhaltens innerhalb einer Betriebsumgebung.

Die Überprüfung der Übertragbarkeit formalisierter Konzepte und Regeln von den analysierten Szenarien auf den offenen Kontext ist nicht Bestandteil des vorgeschlagenen Ansatzes. Grenzen des Ansatzes können daher in Hinblick auf die Repräsentativität des genutzten Szenarienkatalogs und die Validität der verwendeten Wissensquellen identifiziert werden. In Bezug auf die verwendeten Wissensquellen bezieht sich die Validität auf die Akzeptanz der Auswahl aus der Sicht relevanter Stakeholder in der Betriebsumgebung.

5 Anwendung des Ansatzes an einem Beispiel

Zur Demonstration der semantischen Normverhaltensanalyse werden die zwei ausgewählten Szenarien beispielhaft auf Verhaltensnormen analysiert. Hierbei wird als Beispiel einer Wissensquelle ein Abschnitt der Straßenverkehrsordnung (StVO) [31] verwendet. Anschließend wird das Sollverhalten auf Basis des analysierten Normverhaltens formalisiert¹.

¹Die beispielhaft gezeigte Anwendung rechtswissenschaftlicher Methoden soll die mögliche Schnittstelle zwischen der Analyse von Gesetzestexten und der technischen Implementierung eines Regelwerks in der semantischen Normverhaltensanalyse zeigen. Die Übertragbarkeit der umgesetzten Analyse auf andere Wissensquellen, v.a. auf andere Rechtskontexte, wird in dieser Arbeit explizit nicht angenommen. Da die semantische Interpretation von Konzepten und Regeln den Expert*innen der Rechtsdomäne obliegt, macht der Ansatz lediglich einen methodischen Vorschlag zur Anbindung an eine formale Wissensrepräsentation.

Abbildung 3 zeigt zwei funktionale Szenarien (nach [29, 30]), in denen das automatisierte Fahrzeug von Westen in eine T-Kreuzung einfährt und ein Fußgänger von Süden auf einen Fußgängerüberweg der T-Kreuzung zuläuft. In Abbildung 3a ist der südliche Einlaufbereich des Fußgängerüberwegs vollständig durch das automatisierte Fahrzeug einsehbar. Im Gegensatz dazu behindert ein parkendes Fahrzeug in Abbildung 3b die Einsicht in den Einlaufbereich bis zum Haltepunkt am Fußgängerüberweg.

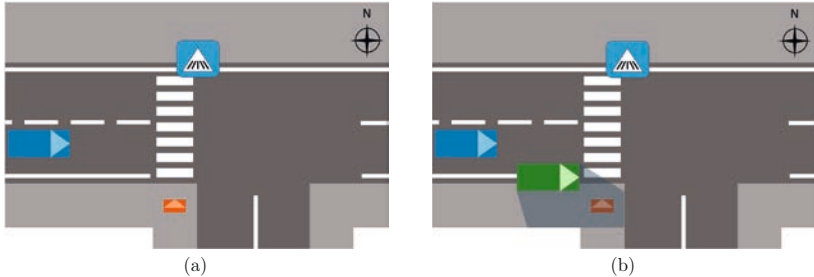


Abbildung 3: Die gewählten funktionalen Szenarien beinhalten das Ego-Fahrzeug (blau), ein parkendes Fahrzeug (grün) und einen auf den Fußgängerüberweg zulaufenden Fußgänger (orange). Der aus Sicht des Ego-Fahrzeugs verdeckte Bereich wurde zur Illustration ergänzt und ist nicht Teil der Szenenbeschreibung.

Der erste Schritt der semantischen Normverhaltensanalyse (Abbildung 2) erfordert eine Definition der zu beachtenden Wissensquellen innerhalb der Betriebsumgebung, für die hier die folgenden Annahmen gelten:

- (A1) Das automatisierte Fahrzeug wird ausschließlich in Deutschland betrieben.
- (A2) Das automatisierte Fahrzeug wird auf öffentlichen Straßen betrieben.
- (A3) Das automatisierte Fahrzeug wird im urbanen Raum betrieben.

5.1 Exemplarische Rechtsanalyse

Als Wissensquelle für die Konzeptualisierung sowie die Formalisierung von Regeln wird beispielhaft ein Abschnitt der deutschen StVO in der Fassung von 2013 [31] verwendet. Weiterhin wird die Allgemeine Verwaltungsvorschrift zur Straßenverkehrsordnung (VwV-StVO) in der Fassung von 2021 [32] herangezogen.

Um Wissen über Konzepte und Regeln, welches in den gewählten Wissensquellen enthalten ist, interpretieren zu können, werden in diesem Beispiel Methoden zur Erstellung rechtswissenschaftlicher Gutachten [33, 34] genutzt. Das bedeutet, dass hier beispielhaft geleistete Interpretationen der StVO² durch den Gutachtenstil und die verwendete

²Es ist zu beachten, dass Regeln für fahrende Personen in diesem Beispiel auf Regeln für das automatisierte Fahrzeug übertragen werden. Die offene Debatte über die juristische Validität dieses Vorgehens wird in diesem Artikel nicht diskutiert (vgl. hierzu Gstöttner u. a. [35]).

Analysemethode maßgeblich die formalisierten Konzepte und Regeln beeinflussen. Gleichzeitig ermöglicht die explizite Interpretation eine durchgängige Dokumentation getroffener Annahmen bei der Formalisierung.

Im **Obersatz** des Gutachtens wird zunächst eine zentrale Fallfrage formuliert, die „neutral und ohne unwichtige Exkurse“ [34, S. 2] im Laufe des Gutachtens beantwortet werden soll. Nach Hildebrand [34] muss das Gutachten dabei vor allem neutral und ökonomisch sein. Es ist also eine Wertung durch den Autor zu vermeiden und es sind ausschließlich die für die Beantwortung der Fallfrage notwendigen Informationen im Gutachten zu repräsentieren. Ausgehend von dem vorliegenden Szenario könnte der Obersatz für das Anwendungsbeispiel folgendermaßen formuliert werden.

Obersatz. *Welche Pflichten entstehen für die fahrende Person des blauen Fahrzeugs nach §26 StVO in der in Abbildung 3a abgebildeten Szene.*

Diese Formulierung des Obersatzes³ ist aufgrund seines konkreten Bezugs nicht ohne Anpassung auf das zweite Szenario in Abbildung 3b anwendbar. Die Evaluation wird zeigen, dass für jedes zu betrachtende Szenario eine semantische Analyse der Verhaltensnormen notwendig ist, um ein korrektes Sollverhalten abzuleiten.

Im nächsten Schritt des gutachterlichen Vierschritts, der **Definition**, werden alle relevanten Rechtsnormen, die zur Beantwortung der Fallfrage notwendig sind, herangezogen und so weit analysiert, wie es zur Beantwortung der Fallfrage nötig ist. Bei der Definition werden alle relevanten Tatbestände und Tatbestandsmerkmale gesammelt und erklärt. Tatbestände werden als Grundlage einer Rechtsfolge verstanden. Für einen Tatbestand sind Merkmale definiert, die für dessen Gültigkeit heranzuziehen sind. Im betrachteten Anwendungsbeispiel könnte ein Teil der Definition lauten:

Definition. *Nach §26 (1) Satz 1 StVO besteht für die fahrende Person eines Fahrzeugs die Pflicht, „an Fußgängerüberwegen den zu Fuß Gehenden, [...], welche den Fußgängerüberweg erkennbar nutzen wollen, [...], das Überqueren der Fahrbahn zu ermöglichen.“ Laut der VwV-StVO zu §26 StVO IV. erfolgt die Kennzeichnung eines Fußgängerüberwegs mit der Markierung Zeichen 293. Zusätzlich wird durch Zeichen 350 auf Fußgängerüberwege hingewiesen.*

In der Definition wurde herausgearbeitet, dass §26 StVO sich unter anderem auf den Tatbestand des Vorhandenseins eines Fußgängerüberwegs bezieht. Die Tatbestandsmerkmale, anhand derer dieser Tatbestand erkannt werden kann, sind in der VwV-StVO zu §26 StVO definiert. Für einen gültigen Fußgängerüberweg ist demnach die Markierung „Zeichen 293“ (der *Zebrastrreifen*) obligatorisch. Zusätzlich weist das blaue Hinweisschild „Zeichen 350“ fahrende Personen auf das Vorhandensein eines Fußgängerüberwegs hin. Das Beispiel zeigt, dass die Erkenntnisse aus der Definition genutzt werden können, um Konzepte und Schlussregeln für die Wissensmodellierung abzuleiten. Nachdem durch den Obersatz eine durch die Analyse zu beantwortende Fragestellung vorgegeben wurde, kann der Definitionsschritt einen Beitrag zur expliziten Interpretation von Gesetzestexten leisten. In der exemplarischen Anwendung wurde die Definition als eine wesentliche Schnittstelle zwischen Expert*innen der Wissensquellen und der Anwendungsdomäne identifiziert.

³Üblicherweise wird der Obersatz eines Gutachtens als Entscheidungsfrage formuliert. Im Beispiel der semantischen Normverhaltensanalyse liegt der Fokus auf der Konzeptualisierung und Regelkonstruktion, die Fragestellung ist daher offen gehalten.

In der **Subsumtion** des Gutachtens werden das Szenario und das terminologische Gerüst, das in der Definition aufgespannt wurde, miteinander verbunden. Die abstrakten Schlüsselbegriffe der Definition werden auf das vorliegende Szenario angewendet (sylogistische Schlussfolgerung), indem verglichen wird, ob die im Szenario vorliegenden Sachverhalte anhand der definierten Tatbestandsmerkmale dem in der Gesetzesnorm vorliegenden Tatbestand zugeordnet werden können [34, S. 26 ff.]):

Subsumtion. *An der in Abbildung 3a abgebildeten Kreuzung ist eine Straßenmarkierung vorhanden, die sich als Zeichen 293 der StVO einordnen lässt. Zudem enthält das Szenario am Straßenrand neben der Markierung 293 das Schild 350 als Hinweis auf einen Fußgängerüberweg. Somit liegt im Szenario, das in 3a abgebildet ist, laut VwV-StVO zu §26 StVO IV. die Kennzeichnung eines Fußgängerüberwegs vor.*

Der letzte Schritt zur Erstellung eines Gutachtens ist das Ausformulieren eines **Ergebnisses**. Das Ergebnis beantwortet die im Obersatz formulierte Frage. Hier wird auf alle Schlussfolgerungen, die in der Subsumtion getroffen wurden, Bezug genommen und ein Gesamtergebnis formuliert:

Ergebnis. *Die fahrende Person des blauen Fahrzeugs aus der in Abbildung 3a beschriebenen Situation ist nach §26 StVO dazu verpflichtet, dem zu Fuß Gehenden das Überqueren der Fahrbahn zu ermöglichen. Somit darf sie nur mit mäßiger Geschwindigkeit heranfahren und muss, wenn nötig, warten.*

Um ein Ergebnis zu formulieren, welches die Frage des Obersatzes abschließend klärt, sind im gewählten Szenario weitere Analysen durchzuführen. Im Anwendungsbeispiel wurden hierfür Annahmen bei der Definition und Subsumtion getroffen. Eine Annahme (basierend auf Annahme A3) innerhalb des Gutachtens war die Markierung des Fußgängerüberwegs innerhalb einer Ortschaft. Diese Annahme wäre im Rahmen einer umfangreicheren Definition zu untersuchen, da sie explizit für die Gültigkeit eines Fußgängerüberwegs erfordert ist. Weiterhin wurde angenommen, dass die Person im Szenario den Überweg erkennbar nutzen möchte. Auch diese Annahme ist in einem aussagekräftigen Gutachten zu untersuchen. Aus Sicht der semantischen Normverhaltensanalyse ist der Ergebnissatz die Zusammenfassung des Sollverhaltens in Bezug auf das vorliegende Szenario und daher dazu geeignet, als eine Indikation für die Richtigkeit des inferierten Verhaltens zu dienen.

5.2 Formalisierung der Konzepte und Regeln

Für eine Durchgängigkeit an der Schnittstelle zwischen natürlich-sprachlicher und formaler Repräsentation des Sollverhaltens ist die Wahl des Abstraktionsgrads wesentlich. Im folgenden Teil des Anwendungsbeispiels wird Sollverhalten auf der Ebene funktionaler Szenarien [29, 30] beispielhaft in eine formale Repräsentation übersetzt.

Eine Entscheidung bei der beispielhaften Umsetzung der semantischen Normverhaltensanalyse wird bei der Wahl der genutzten formalen Repräsentationssprache getroffen. Die Wahl der Sprache zur Repräsentation des Wissens kann durch unterschiedliche sprachliche Mittel Limitationen bei der Beschreibung analysierter Konzepte und Regeln hervorrufen. In diesem Fall müsste die Repräsentationssprache so angepasst werden, dass Wissen aus den Wissensquellen im Sinne der Autor*innen der Wissensquelle abbildbar ist. Im Anwendungsbeispiel werden Konzepte und ihre Beziehungen in der Web Ontology Language (OWL)

beschrieben. Diese Beschreibungslogik eignet sich, um terminologisches Expertenwissen konsistent und maschinenlesbar zu repräsentieren [36].

Eine Eigenschaft der Beschreibungslogik in Bezug auf die Formalisierung von Norm- und Sollverhalten ist, dass sie keine expliziten sprachlichen Mittel zur Repräsentation von Zeit bereitstellt. Diese Eigenschaft wurde adressiert, indem die Konzepte und Regeln bezüglich einer Szene [8] formuliert wurden.

Zusätzlich wurden für die Beschreibung der Szenen des funktionalen Szenarios Zonen [37] verwendet und ohne weitere parametrische Einschränkungen Szenenelemente zugewiesen. Die Zonierung und semantische Normverhaltensanalyse bilden einen iterativen Prozess, der in diesem Beitrag nicht ausgeführt wird und dessen Ergebnis als gegeben betrachtet wird.

Bei der ersten Konzeptualisierung der StVO ergibt sich ein in Abbildung 4 gezeigter taxonomischer Ausschnitt der Klassenstruktur (oder TBox, *terminological box*). Während die abstraktesten vier Klassen expertenbasiert erzeugt werden, ergibt sich die Dekomposition der Klassenstruktur für die Verkehrsinfrastruktur aus der StVO. Das Konzept des Fußgängerüberwegs wird hier als Tatbestand interpretiert, da an ihn die Rechtsfolge im Sinne des analysierten Abschnitts der StVO geknüpft ist. Zusätzlich kann die Klassenstruktur zum Beispiel durch eine domänenspezifische Konzeptualisierung der Betriebsumgebung nach [38–40] vorgegeben sein.

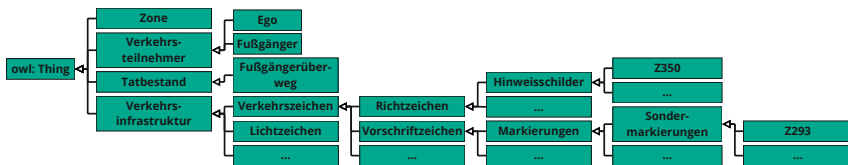


Abbildung 4: In der TBox (*terminological box*) der Ontologie wurde zunächst eine taxonomische Klassenstruktur aus der StVO abgeleitet und zusätzlich um Klassen erweitert, die für die Szenenrepräsentation benötigt werden.

Abbildung 5 zeigt für eine Szene des ersten Szenarios (vgl. Abbildung 3a) die modellierten Instanzen (ABox, *assertional box*) der Klassenstruktur der Ontologie. In der Szene befindet sich das Ego-Fahrzeug im westlichen Straßenast (Zone „blau 1“) und fährt in Richtung Osten auf die Kreuzung zu. Der Fußgängerüberweg befindet sich in Zone „rot“ und im Pfad des Ego-Fahrzeugs. Gemäß der StVO besteht ein Fußgängerüberweg aus einem entsprechenden Hinweisschild (Zeichen 350) und einer Markierung (Zeichen 293). Außerdem befindet sich ein Fußgänger im Einlaufbereich des Fußgängerüberwegs (Zone „grün 1“) und bewegt sich auf ihn zu.

Anschließend folgt die Konstruktion eines Regelkatalogs, der das Sollverhalten basierend auf den gewählten Wissensquellen in den betrachteten Szenarien formalisiert beschreiben soll. Die Konstruktion des Regelkatalogs besteht grundsätzlich aus dem Schritt der Regelableitung und der Regelstrukturierung und ggf. -verfeinerung. Die abgeleiteten Regeln können zunächst in natürlicher Sprache festgehalten werden. Bei der Formalisierung der dekomponierten Regeln wird die Notwendigkeit einer iterativen Methode deutlich: Da Regeln sich teilweise auf implizite Konzepte beziehen, ist es erst bei der Formalisierung der Regel möglich, diese impliziten Konzepte explizit in der Ontologie zu formalisieren. Im Beispiel ist das Konzept „erkennbar queren wollen“ nicht explizit beschrieben. Für

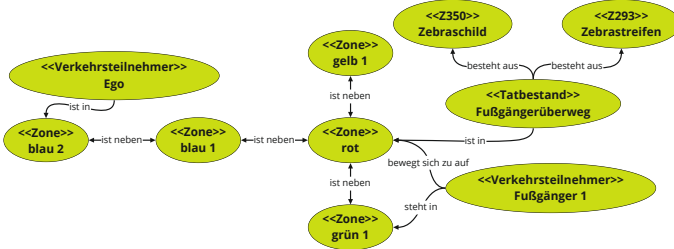


Abbildung 5: In der ABox (*assertional box*) der Ontologie sind Instanzen der Klassen aus der TBox modelliert, die das erste betrachtete Szenario repräsentieren.

die Anwendbarkeit der Regel sind dieses Konzept und die damit verbundenen Annahmen zentral, weshalb es für eine semantische Formalisierung explizit beschrieben werden muss.

In unserem Anwendungsfall wurde § 26 (1) StVO in vier Regeln zerlegt und formalisiert. Eine Hierarchisierung der Regeln wird in diesem Beispiel nicht gezeigt, da sich aus der StVO im Beispiel keine aufzulösenden Konflikte ergeben. Zwischen den Beschreibungsebenen der nicht-formalisierten, natürlich-sprachlichen Wissensquellen und den formalisierten, maschinenlesbaren Regeln kann der Übersetzungsprozess durch folgende semi-formale Regeln unterstützt werden:

- **Wenn** Zeichen 293 *und* Zeichen 350 **dann** gilt Fußgängerüberweg.
- **Wenn** Person in Zone grün 1 (Einlaufzone) **dann** Person will Überweg erkennbar benutzen.
- **Wenn** Fußgängerüberweg in Fahrweg *und* Person will erkennbar benutzen **dann** ist Überqueren zu ermöglichen.
- **Wenn** Regel 1 *und* Regel 2 *und* Regel 3 **dann** Halten in Zone blau 1 (Haltezone vor Fußgängerüberweg).

Die hier vorgestellte Ableitung von Regeln basiert auf dem Wissen, das durch das beispielhafte rechtswissenschaftliche Gutachten gewonnen wurde. Während der Erstellung des Gutachtens wurde iterativ das Vorhandensein von Konzepten und Regeln geprüft und so die Schnittstelle zwischen der Wissensquelle und der formalen Repräsentation hergestellt.

Im Anwendungsbeispiel wird als Notationssprache formaler Regeln die *Semantic Web Rule Language (SWRL)* [10] verwendet. Die implementierten SWRL-Regeln sind im Anhang dieses Artikels (Tabelle 1) aufgeführt.

Die formulierten Regeln werden im folgenden Abschnitt exemplarisch auf die zwei betrachteten Szenarien angewendet und das Ergebnis in Bezug auf das inferierte Sollverhalten auf Regelkonformität hin untersucht.

6 Evaluation der Beispielanwendung

Der letzte Schritt der semantischen Normverhaltensanalyse, die Überprüfung des formalisierten Wissens, sollte in einem Szenarienkatalog erfolgen. Für unser Beispiel erfolgt die

Evaluation des formalisierten Regelsatzes und der daraus resultierenden Schlussfolgerungen anhand des bereits diskutierten Szenarios und einem zusätzlichen Szenario, das eine Verdeckung im Einlaufbereich des Fußgängerüberwegs enthält.

Innerhalb des Inferenzprozesses werden die in SWRL definierten Regeln auf die Entitäten der ABox der OWL-Ontologie angewendet und logische Schlussfolgerungen berechnet. Als Ergebnis wird die Beziehung „anhalten_in“ zwischen dem Ego-Fahrzeug und der im Beispiel als Zone „blau 1“ bezeichneten Zone inferiert. Für die Analyse der Korrektheit der genutzten Regeln soll zukünftig eine Integration des definierten Regelwerks innerhalb eines Phänomen-Signal-Modells [19, 20] erfolgen.

Als Validierung der Übersetzung der Wissensquellen in die vorgestellte Ontologie und in den dazugehörigen Regelkatalog wird das inferierte Verhalten auf Basis von Expertenwissen und der durchgeführten rechtswissenschaftlichen Analyse überprüft. Hierbei fällt auf, dass das Ego-Fahrzeug bei regelkonformem Verhalten immer vor dem Fußgängerüberweg anhalten würde, wenn ein Fußgänger sich im Einlaufbereich befindet. In diesem einfachen funktionalen Szenario zeigt sich zum Beispiel bei der Berücksichtigung des erkennbaren Überquerungswunsches die Herausforderung der Ableitung und Konkretisierung von Regeln. Unter welchen Bedingungen ein Fußgänger tatsächlich als querungswillig gesehen werden kann, kann beispielsweise Gerichtsurteilen, Feldbeobachtungen oder Prädiktionsmodellen [41] entnommen werden. Um in diesem Beispiel nicht weitere Wissenquellen einbeziehen zu müssen, wurde die Annahme getroffen, dass ein im Einlaufbereich stehender Fußgänger erkennbar den Fußgängerüberweg queren will. Diese Annahme wird durch die entsprechende Regel explizit dokumentiert.

Die Analyse des zweiten funktionalen Szenarios zeigt eine weitere Herausforderung bei der Skalierbarkeit des Ansatzes innerhalb eines Szenarienkatalogs. Da das zweite Szenario bei der Konzeptualisierung und Regelkonstruktion nicht explizit berücksichtigt wurde, ist das Konzept der Verdeckung weder in der Ontologie noch in den Regeln berücksichtigt.

Eine automatisierte Prüfung fehlenden Wissens ist bisher nicht Teil des Ansatzes. Daher ist die expertenbasierte Überprüfung aller Szenarien notwendig und es können ausschließlich logische Fehler in der Ontologie und im Regelkatalog automatisiert festgestellt werden.

7 Zusammenfassung und Ausblick

Die semantische Normverhaltensanalyse ist eine Methode zur durchgängigen, formalen Verhaltensspezifikation automatisierter Straßenfahrzeuge. Es wurde argumentiert, dass eine explizite Repräsentation des Norm- und Sollverhaltens einen Beitrag zur Absicherung automatisierter Straßenfahrzeuge in Bezug auf die Erklärbarkeit leisten kann.

Die vorgestellte Analyse ist als Fallbeispiel und ist nicht als rechtssichere Auslegung der StVO zu verstehen. Im Rahmen der Arbeit wurde ein methodischer Vorschlag erarbeitet, um Expert*innen verschiedener Domänen eine Schnittstelle bereitzustellen, damit Interpretationen von Wissensquellen wie der StVO sinngemäß, formal repräsentiert werden können.

Ein Vorteil des Ansatzes ist, dass die Bewertung des inferierten Sollverhaltens durch Expert*innen explizit in den iterativen Entwicklungsprozess einbezogen werden kann. Da die szenarienbezogene Bewertung als Stärke von Analysen durch den Menschen verstanden wird, kann der Ansatz zur semantischen Normverhaltensanalyse als Unterstützungsmittel zur Herleitung von Verhaltensregeln gesehen werden. Die Stärke der formalen Repräsentati-

on der Verhaltensregeln lässt sich insbesondere unter Berücksichtigung eines umfangreichen Szenarienkatalogs und der damit verbundenen Überprüfung von Widersprüchen im Sollverhalten prognostizieren.

Grenzen des Ansatzes ergeben sich beim Übergang zwischen der nicht-formalen und formalen Beschreibung von Verhalten. Die Interpretation von Verhaltensnormen und die Ableitung von Sollverhalten stellt – insbesondere bei der Auflösung von konfliktärem Wissen – eine Fehlerquelle dar. Die explizite Dokumentation von Interpretation an diesem Übergang ist wesentlich für die Argumentierbarkeit getroffener Entwurfsentscheidungen. Eine Herausforderung des offenen Kontextes, die der Ansatz nur eingeschränkt adressiert, ist die Vereinfachung der Realität bei der Verwendung eines szenarienbasierten Ansatzes. Die Formulierung von Verhaltensnormen kann zur Argumentierbarkeit von Annahmen in beobachteten kritischen Grenzfällen beitragen.

Zukünftige Forschung in diesem Kontext könnte die Anbindung semi-formaler Sprachen (vgl. [13]) untersuchen. Auch die Untersuchung weiterer Rechtsquellen und Rechtskontexte für die Anwendung der semantischen Normverhaltensanalyse stellt eine offene Herausforderung dar.

Danksagung

Diese Forschungsarbeiten wurden zum Teil im Rahmen des Projekts „VVMethoden“ durchgeführt. Wir bedanken uns für die finanzielle Unterstützung des Projekts durch das Bundesministerium für Wirtschaft und Klimaschutz (BMWK). Wir bedanken uns insbesondere bei Christian Lalitsch-Schneider (ZF Friedrichshafen AG) und Dr.-Ing. Christoph Höhmann (Mercedes-Benz Group AG) für die anregenden Diskussionen, sowie bei Prof. Dr. rer. nat. Markus Brandstätter (PROSTEP AG) für viel hilfreiches Feedback. Für die Übernahme des Lektorats danken wir Kim Steinkirchner (PROSTEP AG). Abschließend bedanken wir uns bei Hans Nikolaus Beck (ehem. Robert Bosch GmbH), der diese Arbeit fundamental geformt und begleitet hat.

Literatur

- [1] „Taxonomy and Definitions for Terms Related to On-Road Motor Vehicle Automated Driving Systems,“ Society of Automotive Engineers International, Geneva, Switzerland, Technical Report SAE J3016, 2021.
- [2] „Road vehicles – Safety of the intended functionality,“ International Organization for Standardization, Geneva, Switzerland, International Standard ISO 21448, 2022.
- [3] M. Nolte, S. Ernst, J. Richelmann und M. Maurer, „Representing the Unknown – Impact of Uncertainty on the Interaction between Decision Making and Trajectory Generation,“ in *2018 21st International Conference on Intelligent Transportation Systems (ITSC)*, Maui, Hawaii, USA, 2018, S. 2412–2418. DOI: 10.1109/ITSC.2018.8569490.
- [4] A. Censi, K. Slutsky, T. Wongpiromsarn, D. Yershov, S. Pendleton, J. Fu und E. Frazzoli, „Liability, Ethics, and Culture-Aware Behavior Specification using Rulebooks,“ in *2019 International Conference on Robotics and Automation (ICRA)*, ISSN: 2577-087X, 2019, S. 8536–8542. DOI: 10.1109/ICRA.2019.8794364.

- [5] R. C. Arkin, *Behavior-Based Robotics*, 1. Aufl. Cambridge, MA, USA: MIT Press, 1998.
- [6] M. J. Mataric und F. Michaud, „Behavior-based systems,“ in *Springer Handbook of Robotics*, B. Siciliano und O. Khatib, Hrsg., Berlin, Heidelberg: Springer, 2008, S. 891–909. DOI: 10.1007/978-3-540-30301-5_39.
- [7] D. D. Walden, G. J. Roedler, K. Forsberg, R. D. Hamelin und T. M. Shortell, Hrsg., *Systems Engineering Handbook: A Guide for System Life Cycle Processes and Activities*, 4. Aufl., Hoboken: Wiley, 2015.
- [8] S. Ulbrich, T. Menzel, A. Reschka, F. Schuldt und M. Maurer, „Defining and Substantiating the Terms Scene, Situation, and Scenario for Automated Driving,“ in *2015 IEEE 18th International Conference on Intelligent Transportation Systems*, Gran Canaria, Spain, 2015, S. 982–988. DOI: 10.1109/ITSC.2015.164.
- [9] R. Graubohm, T. Stolte, G. Bagschik und M. Maurer, „Towards Efficient Hazard Identification in the Concept Phase of Driverless Vehicle Development,“ *2020 IEEE Intelligent Vehicles Symposium (IV)*, S. 1297–1304, 2020. DOI: 10.1109/IV47402.2020.9304780.
- [10] I. Horrocks, P. F. Patel-Schneider, H. Boley, S. Tabet, B. Grosz und M. Dean, „SWRL: A semantic web rule language combining OWL and RuleML,“ W3C Member submission, 2004.
- [11] P. Feth, R. Adler, T. Fukuda, T. Ishigooka, S. Otsuka, D. Schneider, D. Uecker und K. Yoshimura, „Multi-aspect safety engineering for highly automated driving: Looking beyond functional safety and established standards and methodologies,“ in *Developments in Language Theory*, M. Hoshi und S. Seki, Hrsg., Bd. 11088, Cham: Springer International Publishing, 2018, S. 59–72. DOI: 10.1007/978-3-319-99130-6_5.
- [12] J. Bach, S. Otten und E. Sax, „Model based scenario specification for development and test of automated driving functions,“ in *2016 IEEE Intelligent Vehicles Symposium (IV)*, 2016, S. 1149–1155. DOI: 10.1109/IVS.2016.7535534.
- [13] P. Irvine, X. Zhang, S. Khastgir und P. Jennings, „Structured natural language for expressing rules of the road for automated driving systems,“ in *2023 IEEE Intelligent Vehicles Symposium (IV)*, 2023.
- [14] F. Glatzki, M. Lippert und H. Winner, „Behavioral Attributes for a Behavior-Semantic Scenery Description (BSSD) for the Development of Automated Driving Functions,“ in *2021 IEEE International Intelligent Transportation Systems Conference (ITSC)*, 2021, S. 667–672. DOI: 10.1109/ITSC48978.2021.9564892.
- [15] F. Glatzki und H. Winner, „Inferenz von Verhaltensattributen der Verhaltenssemantischen Szeneriebeschreibung für die Entwicklung automatisierter Fahrfunktionen,“ in *14. Workshop Fahrerassistenzsysteme und automatisiertes Fahren*, Berkheim, Germany, 2022.
- [16] M. Lippert, F. Glatzki und H. Winner, „Behavior-Semantic Scenery Description (BSSD) of Road Networks for Automated Driving,“ *arXiv:2202.05211 [cs, eess]*, 2022.

- [17] G. Bagschik, T. Menzel, C. Körner und M. Maurer, „Wissensbasierte Szenariengenerierung für Betriebsszenarien auf deutschen Autobahnen,“ in *12. Workshop Fahrerassistenzsysteme und automatisiertes Fahren*, Walting im Altmühltal, 2018, S. 1–14.
- [18] M. Butz, C. Heinzemann, M. Herrmann, J. Oehlerking, M. Rittel, N. Schalm und D. Ziegenbein, „SOCA: Domain Analysis for Highly Automated Driving Systems,“ in *2020 IEEE 23rd International Conference on Intelligent Transportation Systems (ITSC)*, 2020. DOI: 10.1109/ITSC45102.2020.9294438.
- [19] H. N. Beck, N. F. Salem, V. Haber, M. Rauschenbach und J. Reich, *Phänomen-Signal-Modell: Formalismus, Graph und Anwendung*, Number: arXiv:2108.00252, 2021. DOI: 10.48550/arXiv.2108.00252.
- [20] —, *Phenomenon-Signal Model: Formalisation, Graph and Application*, Number: arXiv:2207.09996, 2022. DOI: 10.48550/arXiv.2207.09996.
- [21] A. Rizaldi und M. Althoff, „Formalising Traffic Rules for Accountability of Autonomous Vehicles,“ in *2015 IEEE 18th International Conference on Intelligent Transportation Systems*, 2015, S. 1658–1665. DOI: 10.1109/ITSC.2015.269.
- [22] A. Rizaldi, J. Keinholz, M. Huber, J. Feldle, F. Immmler, M. Althoff, E. Hilgendorf und T. Nipkow, „Formalising and Monitoring Traffic Rules for Autonomous Vehicles in Isabelle/HOL,“ übers. von N. Polikarpova und S. Schneider, Ser. Integrated Formal Methods, Cham: Springer International Publishing, 2017, S. 50–66.
- [23] D. Nikol und M. Althoff, „Die Formalisierung von Rechtsnormen am Beispiel des Überholvorgangs,“ *InTeR, Zeitschrift zum Innovations- und Technikrecht*, S. 6, 2019.
- [24] S. Maierhofer, A.-K. Rettinger, E. C. Mayer und M. Althoff, „Formalization of Interstate Traffic Rules in Temporal Logic,“ in *2020 IEEE Intelligent Vehicles Symposium (IV)*, 2020, S. 752–759. DOI: 10.1109/IV47402.2020.9304549.
- [25] D. Hannah und S. Khashtgir, „Proposal for an Approach to Defining Rules of the Road: United Kingdom Proposal,“ UNECE FRAV 17th Session, 2021.
- [26] C. Torens, U. Durak, F. Nikodem und S. Schirmer, „Formally bounding UAS behavior to concept of operation with operation-specific scenario description language,“ in *AIAA Scitech 2019 Forum*, San Diego, California: American Institute of Aeronautics and Astronautics, 2019. DOI: 10.2514/6.2019-1975.
- [27] M. Colledanchise und P. Ögren, „Behavior trees in robotics and AI: An introduction,“ *arXiv:1709.00084 [cs]*, 2018. DOI: 10.1201/9780429489105.
- [28] H.-H. Nagel und M. Arens, „’Innervation des Automobils’ und Formale Logik,“ in *Fahrerassistenzsysteme mit maschineller Wahrnehmung*, Springer, 2005, S. 89–116.
- [29] T. Menzel, G. Bagschik und M. Maurer, „Scenarios for Development, Test and Validation of Automated Vehicles,“ in *2018 IEEE Intelligent Vehicles Symposium (IV)*, 2018, S. 1821–1827. DOI: 10.1109/IVS.2018.8500406.
- [30] G. Bagschik, T. Menzel, A. Reschka und M. Maurer, „Szenarien für Entwicklung, Absicherung und Test von automatisierten Fahrzeugen,“ in *11. Workshop Fahrerassistenzsysteme*, Walting im Altmühltal, 2017, S. 125–135.
- [31] „Straßenverkehrs-Ordnung (StVO),“ Bonn, Verordnung BGB1. I, 2013, S. 367–427.

Bedingte Verhaltensprädiktion interagierender Agenten

Florian Wirth*, Carlos Fernandez-Lopez* und Christoph Stiller*

Zusammenfassung: In der Literatur und in zahlreichen etablierten Benchmarks ist es üblich die zukünftige Position eines anderen Verkehrsteilnehmers (VT) anhand einer Vorhersageverteilung *pro Agenten* zu prädictieren und anhand der Grundwahrheit (GT) dieses Agenten zu evaluieren. Hierfür wird implizit entweder statistische Unabhängigkeit der zukünftigen Positionen der Agenten angenommen oder es wird davon ausgegangen, dass Randverteilungen für die Weiterverarbeitung von Prädiktionen genügen.

In diesem Beitrag wird vorgeschlagen paarweise multivariate Verteilungen als zusätzliche Ausgabe in Vorhersagemodelle für interagierende Agenten zu integrieren. Diese liefern mehr Information falls die Annahme statistischer Unabhängigkeit unzutreffend ist, und andernfalls gleich viel. Aus der paarweise multivariaten Verteilung können dann bedingte Verteilungen gegeben hypothetischer Positionen eines der beiden Agenten – bspw. des automatisierten Ego-Fahrzeugs – abgeleitet werden, wodurch die prädictierte Szene bzgl. verschiedener hypothetischer Aktionen und Aktionskombinationen analysiert werden kann. Der Vorschlag ist für viele Prädiktionsansätze einfach umsetzbar, wenn für jedes Agentenpaar in der Szene eine multivariate Verteilung generiert wird.¹

Schlüsselwörter: Verkehrsszenenprädiktion, bedingte Prädiktion, Verhaltensgenerierung

1 Einleitung

Solange (teil-)manuellgesteuerte Fahrzeuge und Fahrzeuge ohne V2V-Kommunikation im Straßenverkehr teilnehmen, die ihre geplante Trajektorie oder ihre Verhaltensabsicht nicht automatisiert mit anderen Verkehrsteilnehmern teilen können, benötigt ein automatisiertes Fahrzeug (AV) eine Form der Verhaltensvorhersage für die Absichten anderer VT, um sicheres Verhalten zu generieren bzw. eine Trajektorie zu planen. Verhalten zu prädictieren ist also kein Selbstzweck, sondern dient den nachgeschalteten Forschungsfeldern Verhaltensgenerierung und Trajektorienplanung.

Es stellt sich die Frage was eine Prädiktion leisten soll. Grundsätzlich möchte jeder VT schnell, komfortable, ressourcensparend und sicher ein Ziel erreichen. Um auf dem Weg zu diesen Zielen Positionskonflikte zwischen VT effizient zu lösen, interagieren VT unter Zuhilfenahme von Verkehrsregeln. Es bleiben oft mehrere Verhaltensoptionen, wodurch

*Institut für Mess- und Regelungstechnik (MRT), Karlsruher Institut für Technologie (KIT) (E-Mail: florian.wirth@kit.edu).

¹Diese Arbeit basiert auf der Dissertation von Florian Wirth, die am 14.04.2023 bei der Fakultät für Maschinenbau des KIT eingereicht wurde. Details insb. bzgl. des konkreten Netzaufbaus aus dem erwähnten Anwendungsbeispiel in Abschnitt 3.1, der nicht im Fokus des vorliegenden Artikels steht, können dort nachgeschlagen werden.

Verhalten eines andere VT vom eigenen Verhalten abhängen kann oder umgekehrt. Häufig bedingen sich die Verhaltensoptionen zweier oder mehrerer VT gegenseitig und bilden Verhaltensmoden.

Daher genügt es oft nicht zu wissen wo ein anderer VT losgelöst von seinem Umfeld in ein paar Sekunden sein könnte. Eine bedingte Prädiktion erlaubt eine Analyse möglicher Szenenentwicklungen, die passendes Verhalten eines VTs gegeben eines bestimmten Verhaltens eines anderen VTs – bspw. des automatisierten Ego-Fahrzeugs² – liefert.

2 Stand der Technik

In der Literatur können verschiedene Formen der Modellausgabe gefunden werden, die hier aufgezählt und beschrieben werden.

Unabhängige Prädiktion von Agenten. Die meisten konkurrierenden Ansätze geben agentenzentrierte Positionsprädiktionen aus, die nicht zueinander in Bezug gesetzt werden können. Implizit wird Prädiktion hier als Zufallsexperiment modelliert, bei dem die zukünftigen Positionen mehrerer Agenten als statistisch unabhängig angenommen werden. Es kann daraus nicht abgeleitet werden welche Trajektorie von Agent *a* zu welcher Trajektorie von Agent *b* gehört [1, 2, 3, 4, 5, 6].

Obwohl auch sie agentenweise Positionsprädiktionen ausgeben, weisen Khandelwal and Qi *et al.* [7] auf die Notwendigkeit hin hypothetische Szenarienentwicklungen analysieren zu können und evaluieren ihr Prädiktionsmodell mit künstlich eingefügten Agenten, mit denen das Ego-Fahrzeug interagieren müsste, wenn sie tatsächlich existierten.

Ego-Trajektorie als Eingabe. Salzmann, Ivanovic *et al.* [8] fügen die geplante Trajektorie des Ego-Fahrzeugs in deren Modellinput hinzu. Dadurch hat das Modell die Möglichkeit Verhalten anderer Verkehrsteilnehmer bedingt auf verschiedene Ego-Trajektorien zu generieren. Auch Tolstaya *et al.* [9] schlagen ein Modell vor, das sie *Conditional Behavior Prediction* nennen und das von einem Agentenpaar ausgeht bestehend aus Mensch und Roboter, bei dem die Trajektorie des Roboters festgelegt wird und die zugehörige Trajektorie des Menschen prädictiert wird.

Konsistente Verhaltensmoden. In letzter Zeit geht der Trend hin zur Prädiktion konsistenter Verkehrsszenenentwicklungen, die *Moden (modes)* [10, 11] genannt werden. Eine Mode umfasst eine Trajektorie pro prädictiertem Agenten und wenn jeder Agent diese Trajektorie fährt entwickelt sich die Verkehrsszene unfallfrei. Im Allgemeinen gibt es für eine Verkehrsszene mehrere mögliche Moden.

Gilles *et al.* [12] schlagen ein Modell vor, das eine Belegtheitsrasterkarte für jeden Agenten generiert, woraus eine konstante Zahl von Trajektorienendpunkten gezogen wird. Ein nachgeschaltetes Modell lernt aus diesen Endpunkten konsistente Verhaltensmoden zu generieren und jeder Mode eine Auftrittswahrscheinlichkeit zuzuordnen.

Bedingte Prädiktion. Rhinehart *et al.* [13] erheben den Anspruch das erste bedingte Vorhersagemodell für das automatische Fahren entwickelt zu haben. Sie entwickelten ein generatives Modell für einzelne Agenten basierend auf Rhinehart *et al.* [14], das das Verhalten mehrerer Agenten prädictiert, gegeben ihre früheren Trajektorien und einer Rasterkarte. Das tun sie, indem sie das Verhalten aller Agenten mit einem gemeinsamen gene-

²Die gängige Abkürzung “AV” für *autonomous vehicle* wird im Folgenden synonym zum automatisierten Ego-Fahrzeug verwendet.

rativen Modell präzisieren, dessen Eingabe nicht nur die vergangene Trajektorie und die Rasterkarte ist, sondern auch ein Zufallszustand einer Gauß-Verteilung für jeden Agenten.

Wenn die Zuständen, die die gewünschte zukünftige Trajektorie des Roboters repräsentieren, gefunden werden, können die Zustände aller anderen Agenten mit Gauß'schem Rauschen gefüllt werden, während die Modelleingabe für den Roboter mit dem gefundenen Zustand ersetzt wird. Die Autoren formulieren ein Optimierungsproblem, mit dem der gewünschte Zustand für den Roboter generiert werden kann.

Tang *et al.* [15] veröffentlichen ein generatives Modell, das in der Lage ist verschiedene zukünftige Trajektorien *auszurollen* bei beliebig vielen vorgegebenen Wunschtrajektorien. Ein gemeinsames rekurrentes neuronales Netz für alle Agenten generiert zukünftige Positionen für den jeweils nächsten Zeitschritt. Da die Ausgabeposition wieder in das Modell eingespeist wird, kann diese Position für beliebig viele Agenten auf einen vorgegebenen Wert gesetzt werden.

Eine Erweiterung von Casas, Culino, Suo *et al.* [10] wird in Cui, Casas, Sadet *et al.* [16] vorgeschlagen, die direkt konsistente Verhaltensmoden ähnlich zu Gilles *et al.* [12] ziehen, und darüber hinaus eine Planungsmethode vorschlagen, die die generierte Ausgabe weiterverarbeiten kann.

Der Ansatz von Ngiam, Caine *et al.* [17] ist in der Lage eine beliebige Menge bedingter Trajektorienpunkte zu verarbeiten und gibt ebenfalls dazu konsistente Verhaltensmoden mit entsprechenden Wahrscheinlichkeiten aus.

3 Vorgeschlagene Prädiktionsausgabe

Aufgrund der Auswahl an Metriken in populären Benchmarks fokussieren sich viele Forscher auf die Positionsprädiktion einzelner Agenten. Häufig wird pro Agent eine Gauß'sche Mischung ausgegeben.

Im Gegensatz dazu wird vorgeschlagen Agenten paarweise zu präzisieren. Statt 2D Gauß-Mischverteilungen können 4D Gauß-Mischverteilungen generiert werden. Diese stehen nun für paarweise Verhaltensmoden zweier Agenten und können entsprechend trainiert und analysiert werden.

Üblicherweise werden innerhalb von Prädiktionsmodellen zukünftige latente Zustände einzelner Agenten $\mathbf{s}_a \forall a \in [1, \dots, n_A]$ trainiert, aus denen dann die Modellausgabe für diesen Agenten decodiert wird. Hierbei ist n_A die Anzahl der Agenten in der Szene. Stattdessen kann durch Kreuzung dieser Zustände eine Informationsrepräsentation $[\mathbf{s}_a, \mathbf{s}_b] \forall (a, b) \in [1, \dots, n_A]^2$ geschaffen werden, die die Zustände zweier Agenten enthält. Nun kann ein – üblicherweise trainierbarer – Decoder φ_{Decod} entworfen werden, der die 4D Ausgabe anhand der veränderlichen Gewichte im Decoder optimiert gegeben der Zustandspaare $[\mathbf{s}_a, \mathbf{s}_b]$. Da Symmetrie der 4D-Verteilung bezüglich des Agentenpaares $\varphi_{\text{Decod}}([\mathbf{s}_a, \mathbf{s}_b]) \stackrel{!}{=} \varphi_{\text{Decod}}([\mathbf{s}_b, \mathbf{s}_a])$ erwartet wird, kann die Eingabe des Decoders z.B. zu $\varphi_{\text{Decod}}(\mathbf{s}_b + \mathbf{s}_a)$ gewählt werden. Somit ist die Ausgabe des Decoders invariant gegenüber der Reihenfolge von a und b .

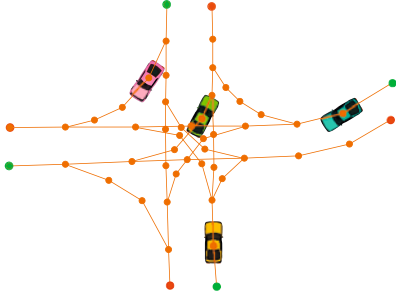


Abbildung 1: Visualisierung der Graphmodellierung: Der Kartengraph ist zerlegt in Fahrbahnstücke, die in etwa die Fläche einnehmen, die ein Auto belegt. Die Position eines Agenten wird festgelegt durch eine diskrete Verteilung über alle Fahrbahnstücke. In dieser Grafik hat die Verteilung jedes Agenten genau einen Eintrag, der eine 1 enthält. Dieser weist dem Agenten genau ein Fahrbahnstück zu.

3.1 Anwendungsbeispiel: Diskrete Prädiktion von Agenten auf einem Kartengraphen

Multivariate Wahrscheinlichkeiten können besonders einfach generiert werden, wenn die Karte als Graph mit diskreten Aufenthaltswahrscheinlichkeiten aufgefasst wird. In Abb. 1 ist ein Kartengraph dargestellt, der aus Fahrstreifenstücken $l \in [1, \dots, n_L]$ besteht, die ungefähr die Größe eines Autos einnehmen. Hierbei ist n_L die Anzahl der Fahrstreifenstücke in der betrachteten Szene. Somit entspricht die maximale Positionsgenauigkeit für einen Agenten der Größe des Agenten, was als ausreichend angenommen wird. Weiterhin wird die aktuelle Position jedes Agenten im Kartengraph bestimmt und zur Prädiktion über n_T diskrete, zukünftige Zeitschritte t werden Zustände der Agenten entlang des Kartengraphen propagiert. Es liegt also im Kartengraph für jeden Agenten a auf jedem Fahrstreifenstück l zu jedem zukünftigen Zeitschritt $t \in [1, \dots, n_T]$ ein latenter, trainierbarer Zustand $\mathbf{s}_{t,l,a}$ vor.

Aus diesem Zustand kann nun wie gehabt eine Aufenthaltswahrscheinlichkeit $p_{t,l,a} = \text{norm}_l(\varphi_{\text{Decod}}(\mathbf{s}_{t,l,a}))$ generiert werden, indem durch Normalisierung über die Fahrstreifenstücke sichergestellt ist, dass die Ausgabe Anforderungen an eine Wahrscheinlichkeitsverteilung genügt. Dies entspricht der Wahrscheinlichkeit $\mathcal{P}$ des Ereignisses Ω , dass sich Agent a zum Zeitpunkt t auf Fahrstreifenstück l befinden wird: $p_{t,l,a} = \mathcal{P}(\Omega_{t,l,a})$. Dem Vorschlag dieses Artikels folgend kann darüber hinaus eine bivariate, diskrete Verteilung für jedes Agentenpaar generiert werden:

$$p_{t,k,b,l,a} = p_{t,l,a,k,b} = \text{norm}_{l,k}(\varphi_{\text{Decod}}(\mathbf{s}_{t,l,a} + \mathbf{s}_{t,k,b})) = \mathcal{P}(\Omega_{t,l,a} \wedge \Omega_{t,k,b}). \quad (1)$$

Die Indizes k, b stehen hierbei für den zweiten Agenten b auf dem Fahrstreifenstück k . Durch Wählen konkreter Werte für b und k (oder a und b) und normalisieren über den jeweils anderen Fahrstreifenstückindex l (oder k) kann eine diskrete bedingte Prädiktion generiert werden für einen Agenten gegeben einer zukünftigen Position eines anderen Agenten. Analog kann über Bereiche einer kontinuierlichen Verteilung integriert und normalisiert werden.

3.2 Analysemöglichkeiten anhand bedingter Vorhersagen

Anhand der Trainingsdaten liefert das Modell das wahrscheinlichste menschliche Verhalten für die vorliegende Verkehrssituation, sofern – wie die meisten heutigen Datensätzen –

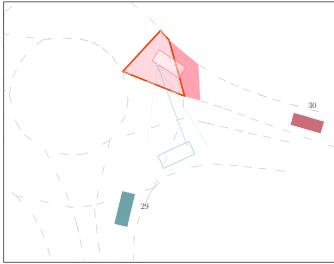


Abbildung 2: Visualisierung der Prädiktion zweier Agenten. Die bedingte Prädiktion wird durch Verbindungslinien in der Farbe des bedingten Agenten (blau) visualisiert. Die Farbintensität repräsentiert hierbei die Wahrscheinlichkeit. Als Bedingung wird das GT Fahrstreifenstück (rot umrandet) des aktuell prädierte Agent verwendet (magenta). Die bedingte Positionsprädiktion für den blauen Agenten gegeben, dass der rote Agent in den Kreisverkehr eingefahren ist, zeigt vorrangig auf eines der Fahrstreifenstücke, die zum ersten Ausgang des Kreisverkehrs gehören.

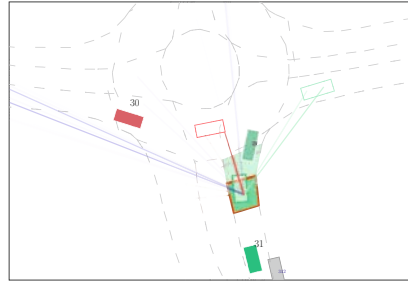
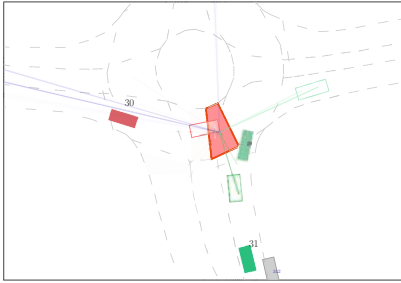
ausschließlich manuell gesteuerte Fahrzeug aufgenommen wurden. Es wird daher empfohlen bedingte Prädiktionen für das AV gegeben der wahrscheinlichsten Position desjenigen Agenten zu generieren, mit dem das AV den größten zukünftigen Positionskonflikt hat. Anhand dieser bedingten Prädiktion kann das AV ein Verhalten oder eine Trajektorie generieren, die menschlichem Fahrverhalten nahekommt. Positionskonflikte können anhand der Überlappung naiver Positionsprädiktionsmethoden bspw. anhand von 1D Kalman-Filtern mit einem konstante-Beschleunigung-Modell entlang der Fahrstreifen generiert werden. Da häufig nur ein Positionskonflikt mit einem anderen Agenten vorliegt, genügt eine paarweise Analyse der Prädiktion von Agenten oft, um Verhalten für kurze Planungshorizonte von wenigen Sekunden zu generieren. Wenn die paarweise Analyse mehrerer anderer Fahrzeuge bzgl. des AV vergleichbares Verhalten für das AV ergeben, wird die paarweise Betrachtung ebenfalls als ausreichend angesehen.

Allerdings sind Beispiele denkbar, bei denen Konflikte nicht nur paarweise auftreten, sondern die gerichtete Kausalkette drei oder mehr Agenten umfassen. Statistische Modelle liefern nur statistische Zusammenhänge, nicht jedoch die Richtung der Kausalbeziehung. Für mehr als zwei Agenten a, b, c , bspw. in einem Rechts-vor-Links-Szenario bei der *kein* Deadlock entsteht, kann überprüft werden, ob das Verhalten für a aus den Paaren (a, b) und (b, c) mit dem Verhalten aus (a, c) übereinstimmt. Der Vorschlag dieses Artikels multivariate Verteilungen für Agentenpaare zu generieren enthält im Gegensatz zu Arbeiten, die konsistente Verhaltensmoden generieren, keinen Mechanismus, der einen solchen statistischen Ringschluss sicherstellt. Der Ringschluss kann also beim vorliegenden Vorschlag dazu verwendet werden, um die Kohärenz von mehr als zwei bedingten Prädiktionen zu überprüfen. Sollte sich der Ringschluss als grob unzulässig erweisen, ist die Prädiktion mit Vorsicht zu behandeln. Die Verkehrssituation muss sich dann ggf. erst mit konservativem Verhalten des AVs weiterentwickeln, um zu einem validen Prädiktionsergebnis zu führen.

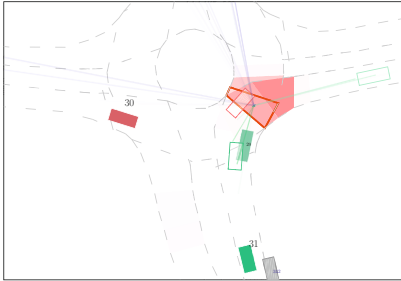
4 Beispiele für die Situationsanalyse

In Abb. 2 wird an einem Beispiel aus dem INTERACTION Datensatz die verwendete Visualisierung für die diskrete Positionsprädiktion und die diskrete bedingte Prädiktion in Kürze erläutert.

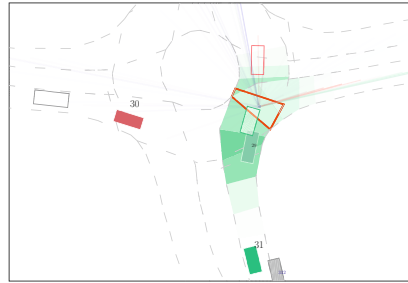
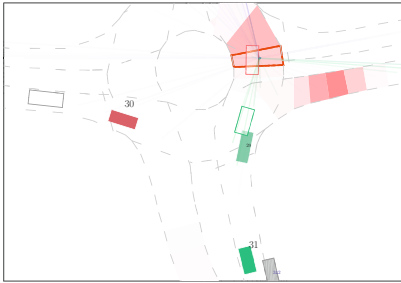
Der Eingang des Modells – die aktuellen Beobachtungen der Agenten – sind als far-



- (a) Die Prädiktionen von Agent #30 (rot) und #31 (grün unten) passen nach 2.1 s gut zur GT. Auch die bedingten Prädiktionen zeigen von der GT zu den Fahrbahnstücken, auf denen der jeweils andere Agent in der Zukunft ist.



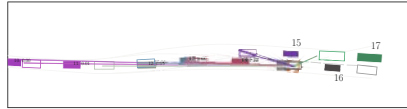
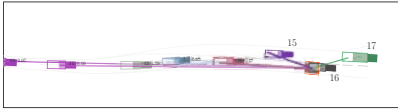
- (b) Auch nach 3.3s zeigen die farbigen Linien zu den farbigen Boxen des jeweils anderen Agenten.



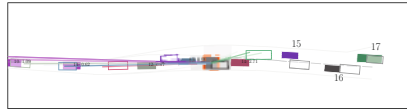
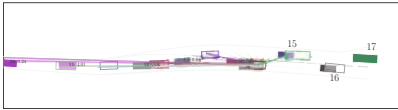
- (c) Für 4.5s macht die große Unsicherheit in longitudinaler Richtung zwischen den Verhaltensoptionen "Bremsen" und "Einfahren" von Agent #31 die Prädiktion weitestgehend unbrauchbar. Die bedingten Prädiktionen passen jedoch für beide Agenten weiterhin gut zur GT des jeweils anderen Agenten.

Abbildung 3: Bedingte Prädiktion, Beispiel 1: roundD-Datensatz, location1.

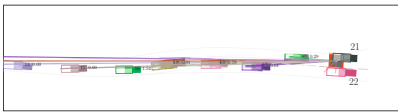
biges Rechtecke mit schwarzer Agenten-ID dargestellt, die zugehörige zukünftige Position als transparente Rechtecke mit farbigem Rahmen. Die Straßentopologie wird in grau angedeutet, es handelt sich um einen Kreisverkehr. Die prädiizierte Belegtheitswahrschein-



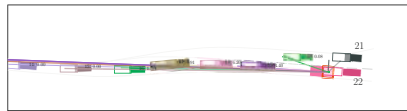
- (a) Links: 0.3 s. Rechts: 1.2 s. Ein Reißverschlusszenario, in welchem der lila Agent #15 vor und der grüne Agent #17 vor dem braunen Agenten #16 einfädelt.



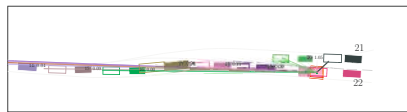
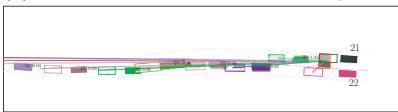
- (b) Links: 2.1 s. Rechts: 3.3 s. Die bedingte Prädiktion dargestellt anhand von lila und grünen Linien geben die Reihenfolge des Einfädelns für kurze und lange Prädiktionshorizonte korrekt an.



- (c) 0.3 s: Einige Sekunden später fahren der pinke Agent #22 und der dunkelgrüne Agent #21 direkt nebeneinander in die Szene ein. Bereits für kurze Prädiktionshorizonte...



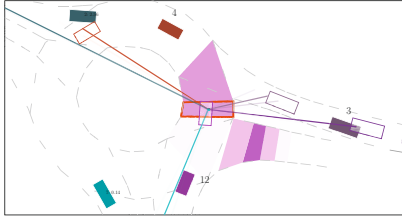
- (d) 0.6 s: ...wird die Einfädelreihenfolge korrekt prädiziert,...



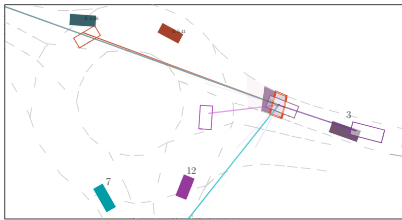
- (e) 0.9 s: ...die vermutlich aus der Geschwindigkeitsdifferenz beider Agenten abgeleitet wurde: Der pinke Agent hat bereits nach der ersten Sekunde eine größere Strecke zurückgelegt, daher wird er bereits nach kurzer Zeit als vor Agent #21 einfädelnd prädiziert.

Abbildung 4: Bedingte Prädiktion, Beispiel 2: INTERACTION-Datensatz, DR_DEU_Merging_MT. Zwei Fahrstreifen verengen zu einer. (a) und (b) zeigen vier verschiedene Prädiktionshorizonte einer ersten Reißverschlusszene. (c), (d) und (e) zeigen zwei Agenten #21 und #22 zu drei verschiedenen Prädiktionshorizonten einer zweiten Reißverschlusszene. In allen gezeigten Beispielen deutet die bedingte Prädiktion mittels farbliger Linien die tatsächlich gefahrene Einfädelordnung korrekt an.

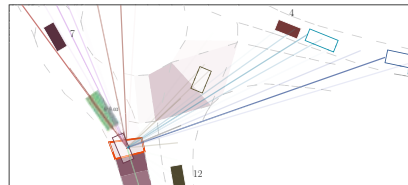
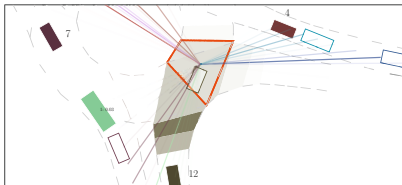
lichkeit wird durch Einfärbung der entsprechenden Fahrstreifenstücke mit der Farbe des zugehörigen Agenten dargestellt (hier magenta für Agent #30). Deren Farbintensität repräsentiert die Wahrscheinlichkeit. Das GT-Fahrbahnstück ist rot umrandet. Die bedingte Prädiktion gegeben des GT-Fahrbahnstücks (rot) des aktuell prädizierten Agenten (magenta) wird durch farbige Linien (blau) visualisiert, die zu den wahrscheinlichen, zugehörigen Fahrbahnstreifen desjenigen Agenten zeigen, dessen Farbe die Linien haben.



- (a) Der lila Agent #12 fährt in den Kreisverkehr und wird mit je ca. 50% wird er den Kreisverkehr an der ersten Ausfahrt verlassen oder darin bleiben. Die wahrscheinlichste Positionsprädiktion für das graue Fahrzeug #3 bedingt darauf, dass Agent #12 im Kreisverkehr bleibt, liegt bei der Stoplinie der Zufahrt: Dies wird durch die graue Linie vom roten GT-Fahrstreifenabschnitt von #12 hin zur Stoplinie der Zufahrt veranschaulicht.



- (b) Umgekehrt wird die bedingte Wahrscheinlichkeit für #12 dargestellt gegeben, dass #3 vor der Stoplinie bremst. Links, 2.1 s: Anfangs ist die Wahrscheinlichkeitsmasse der bedingten Prädiktion mehr oder weniger gleichverteilt zwischen den zwei Verhaltensoptionen von #12 *im Kreisverkehr bleiben* und *die erste Ausfahrt nehmen*. Rechts, 3.3 s: Da #3 bremsen muss, um hinter der Stoplinie zu bleiben und #12 vorbeifahren zu lassen, zeigen die lila Linien nun nur noch in den Kreis des Kreisverkehrs, wodurch die zwei korrespondierenden Verhaltensmoden von #3 und #12 verbunden werden.



- (c) Im Gegensatz dazu verlässt der braune Agent #7 den Kreisverkehr und der grüne Agent #12 kann direkt einfahren. Während die bedingte Prädiktion für #7 gegeben des Einfahrens von #12 gut passt (links, 3.9 s), scheint die bedingte Prädiktion für #12 gegeben, dass #7 den Kreisverkehr verlässt, mehrdeutig (rechts, 3.9 s). Dies könnte von der asymmetrischen Abhängigkeit herrühren: Das Verhalten von #12 hängt von Agent #7 ab, da #12 Vorfahrt gewährend müsste, aber auch von jedem Agenten hinter #7. Agent #7 kann sich unabhängig von #12 bewegen, da er Vorfahrt hat, und unabhängig von jedem Agenten hinter ihm selbst.

Abbildung 5: Beispiel 3: INTERACTION-Datensatz, DR_DEU_Roundabout_OF. Eine dichte Verkehrsszene an einem Kreisverkehr ist abgebildet für verschiedene Agenten und verschiedene Prädiktionshorizonte.

Das zugehörige Verhalten des blauen Agenten #29 gegeben, dass der Agent in magenta #30 in den Kreisverkehr einfährt (rot umrandetes Fahrstreifenstück), ist also ein Verlassen des Kreisverkehrs über die erste Ausfahrt. Die Kausalität aber ist in diesem Beispiel umgekehrt: *Weil* der blaue Agent – der Vorfahrt hätte – die erste Ausfahrt nimmt, fährt der Agent in magenta direkt in den Kreisverkehr.

In Abb. 3, Abb. 4 und Abb. 5 werden anhand dreier Beispiele die Vorzüge und Eigenschaften bedingter Prädiktionen erläutert.

5 Schlussfolgerung

Bedingte Positionsprädiktionen liefern einen deutlichen Mehrwert gegenüber reinen Positionsprädiktionen, bei denen implizit statistische Unabhängigkeit der zukünftigen Positionen von Agenten angenommen wird. In diesem Beitrag wird eine Methode vorgeschlagen, mit der sich bedingte Prädiktionen in viele Prädiktionsverfahren einfach integrieren lassen. Der Vorschlag umfasst die Prädiktion von multivariaten Verteilungen für jedes Agentenpaar einer Verkehrsszene, aus der bedingte Prädiktionen abgeleitet werden können. Hiermit kann eine Vielzahl hypothetischer Entwicklungen einer Verkehrsszene analysiert und für die Verhaltensgenerierung oder die Trajektorienplanung verwendet werden.

Literatur

- [1] W. Zeng, M. Liang, R. Liao, and R. Urtasun, “Lanercnn: Distributed representations for graph-centric motion forecasting,” *2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pp. 532–539, 2021.
- [2] T. Gilles, S. Sabatini, D. Tsishkou, B. Stanciulescu, and F. Moutarde, “Home: Heat-map output for future motion estimation,” in *IEEE International Intelligent Transportation Systems Conference (ITSC)*, (Indianapolis, IN, USA), pp. 500–507, 2021.
- [3] M. Ye, T. Cao, and Q. Chen, “Tpcn: Temporal point cloud networks for motion forecasting,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, (Nashville, TN, USA), pp. 11318–11327, 2021.
- [4] Y. Yuan, X. Weng, Y. Ou, and K. M. Kitani, “Agentformer: Agent-aware transformers for socio-temporal multi-agent forecasting,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, (Montreal, Canada), pp. 9813–9823, October 2021.
- [5] T. Gilles, S. Sabatini, D. Tsishkou, B. Stanciulescu, and F. Moutarde, “Gohome: Graph-oriented heatmap output for future motion estimation,” in *International Conference on Robotics and Automation (ICRA)*, (Philadelphia, PA, USA), pp. 9107–9114, 2022.
- [6] B. Varadarajan, A. Hefny, A. Srivastava, K. S. Refaat, N. Nayakanti, A. Cornman, K. Chen, B. Douillard, C. P. Lam, D. Anguelov, *et al.*, “Multipath++: Efficient information fusion and trajectory aggregation for behavior prediction,” in *International Conference on Robotics and Automation (ICRA)*, (Philadelphia, PA, USA), pp. 7814–7821, 2022.

- [7] S. Khandelwal, W. Qi, J. Singh, A. Hartnett, and D. Ramanan, “What-if motion prediction for autonomous driving,” *CoRR*, vol. abs/2008.10587, 2020.
- [8] T. Salzmann, B. Ivanovic, P. Chakravarty, and M. Pavone, “Trajectron++: Dynamically-feasible trajectory forecasting with heterogeneous data,” in *European Conference on Computer Vision (ECCV)*, (Glasgow, United Kingdom), pp. 683–700, Springer, 2020.
- [9] E. Tolstaya, R. Mahjourian, C. Downey, B. Vadarajan, B. Sapp, and D. Anguelov, “Identifying driver interactions via conditional behavior prediction,” in *2021 IEEE International Conference on Robotics and Automation (ICRA)*, (Xi’an, China), pp. 3473–3479, 2021.
- [10] S. Casas, C. Gulino, S. Suo, K. Luo, R. Liao, and R. Urtasun, “Implicit latent variable model for scene-consistent motion forecasting,” in *European Conference on Computer Vision (ECCV)*, (virtual), pp. 624–641, Springer, 2020.
- [11] S. H. Park, G. Lee, J. Seo, M. Bhat, M. Kang, J. Francis, A. Jadhav, P. P. Liang, and L.-P. Morency, “Diverse and admissible trajectory forecasting through multimodal context understanding,” in *European Conference on Computer Vision (ECCV)*, (Glasgow, United Kingdom), pp. 282–298, Springer, 2020.
- [12] T. Gilles, S. Sabatini, D. Tsishkou, B. Stanciulescu, and F. Moutarde, “Thomas: Trajectory heatmap output with learned multi-agent sampling,” in *International Conference on Learning Representations (ICLR)*, (virtual), 2022.
- [13] N. Rhinehart, R. McAllister, K. Kitani, and S. Levine, “Precog: Prediction conditioned on goals in visual multi-agent settings,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, (Seoul, Korea), October 2019.
- [14] N. Rhinehart, K. M. Kitani, and P. Vernaza, “R2p2: A reparameterized pushforward policy for diverse, precise generative path forecasting,” in *Proceedings of the European Conference on Computer Vision (ECCV)*, (Munich, Germany), pp. 772–788, 2018.
- [15] C. Tang and R. R. Salakhutdinov, “Multiple futures prediction,” *Advances in Neural Information Processing Systems*, vol. 32, 2019.
- [16] A. Cui, S. Casas, A. Sadat, R. Liao, and R. Urtasun, “Lookout: Diverse multi-future prediction and planning for self-driving,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, (Montreal, Canada), pp. 16107–16116, October 2021.
- [17] J. Ngiam, V. Vasudevan, B. Caine, Z. Zhang, H.-T. L. Chiang, J. Ling, R. Roelofs, A. Bewley, C. Liu, A. Venugopal, D. J. Weiss, B. Sapp, Z. Chen, and J. Shlens, “Scene transformer: A unified architecture for predicting future trajectories of multiple agents,” in *International Conference on Learning Representations (ICLR)*, (virtual), 2022.

Systematic Derivation of Use Case Clusters for a Generalized Low-Speed Automated Driving Function

Moritz Berghöfer*, Melina Lutwitz* and Steven Peters*

Abstract: One approach for the introduction of SAE Level 3+ Automated Driving are low-speed driving functions due to a reduced risk associated with them. In this paper, a systematic methodology to derive use cases and use case clusters for low-speed applications of Automated Driving (AD), which can be potentially fulfilled by a generalized low-speed function architecture, is described and applied. The use case clusters are defined according to a classification of the derived use cases in the dimensions of safety and technical capabilities. Thereby, the results of this paper simplify the definition of the ODD as well as the functional requirements and architecture for the future development of low-speed AD functions.

Keywords: Automated Driving, Low-Speed Functions, Use Cases, Requirements, Operational Design Domain

1 Introduction and State of the Art

The field of Automated Driving (AD) has evolved substantially during the last years. A major challenge in this field is the introduction of AD functions with functionality according to SAE Levels 3 and 4 [1] into the existing traffic [2]. One approach to face this challenge, e.g. presented by Bolle et al. [3], is the restriction of use cases for AD. Low-speed AD applications seem very suitable in this context because of a simple fail-safe strategy. Due to the low kinetic energy and the resulting low braking distance the overall risk is highly reduced [2]. Furthermore, low-speed AD applications have minor technical requirements on the perception, prediction, and planning horizon of the AD system [3]. Thereby, low-speed AD functions come along with a reduced effort for the technical development and the safety approval, enabling an earlier market introduction and serving as a basis for introducing AD functions with a more complex Operational Design Domain (ODD). One piece of evidence for that is that the worldwide first regulatory-approved AD function according to SAE level 4, an Automated Valet Parking (AVP) system in a car park at the Stuttgart airport [4], is a low-speed application. Another example that underlines the relevance of speed limitation for the approval of AD functions is the recent introduction of the first internationally approved SAE Level 3 system, the *Drive Pilot* by Mercedes-Benz. Its ODD is limited to highway drives under further conditions, but only to maximum speeds of 60 km/h [5].

*All authors are with the Institute of Automotive Engineering Darmstadt (FZD), TU Darmstadt, Otto-Berndt-Str. 2, 64287 Darmstadt, Germany (e-mail: {moritz.berghoefer; melina.lutwitz; steven.peters}@tu-darmstadt.de).

A second motivation for low-speed AD functions is that they have the potential to increase the comfort and efficiency of daily life in the near future. The aforementioned use case of AVP offers several benefits like the reduction of vehicle damages [6], the increase of time and energy efficiency for searching a parking space [7], and the possibility for High Density Parking (HDP) [6]. AVP is currently the most popular and evolved application for low-speed AD functions, as evidenced by the fact that there are regulatory documents and standards in development and release, as a technical requirements catalog for AVP published by the Kraftfahrt-Bundesamt [8] and two ISO standards under development (ISO 23374-1 [9], ISO 12768 [10]), which focus on AVP and possible extensions. As the scope of ISO 12768 shows, AVP can also be extended to use cases like automated charging (*Automated Valet Charging*, AVC), e.g. [11] and [12], or washing (*Automated Valet Washing*, AVW), e.g. [13]. Beyond that, low-speed shuttle applications are already in the focus of research and industry for longer than 20 years [14]. Low-speed shuttle services seem beneficial especially for industrial applications, where goods, persons, or the vehicle itself (see e.g. pilot project *Automatisiertes Fahren im Werk* (AFW) by BMW [15]) are transported on a fixed route from a starting point to a certain destination.

To summarize, there are a lot of different perspectives and stakeholders, like society, OEM, and vehicle owners, that have an interest in low-speed AD functions. It therefore is desirable to expand the potential of low-speed AD functions by exploiting further use cases, e.g. also in the regular road traffic. Till now there exist only a few approaches to holistically consider use cases and scenarios of low-speed AD functions. At this point, ISO 22737 (“Low-speed automated driving (LSAD) systems for predefined routes”) [16] should be mentioned, which defines overall system, safety, and performance requirements for low-speed applications on predefined routes. However, neither specific use cases nor the important aspect of technical requirements, e.g. the necessary sensor setup, are considered.

In order to close this gap in the field of AD in low-speed applications, in the scope of the public-funded project AUTotech.agil [17] one goal is to develop a generalized AD function for low-speeds (“Generalized Low-Speed Function”, GLSF), which uses a dedicated short-range sensor setup as well as own perception and planning modules. This driving function shall cover as many different use cases and scenarios in the low-speed range as possible with one generic architecture. In the context of a modular platform architecture of automated vehicles, as realized in the project UNICARagil [18], the GLSF as one unique function can thus be implemented in different vehicle concepts (e.g. private-owned vehicle vs. cargo shuttle).

As a starting point for the specification of the function’s requirements and architecture, it is necessary to define the use cases to be implemented and the ODD under consideration. To select the use cases, it is necessary to look for common features in possible use cases for low-speed AD functions. For this purpose, this paper presents a methodology for the systematic derivation of use case clusters for low-speed AD functions. This enables to frame the scope of a generalized low-speed AD function and its required capabilities, which will in the future be used to derive its architecture as well as its functional and technical requirements.

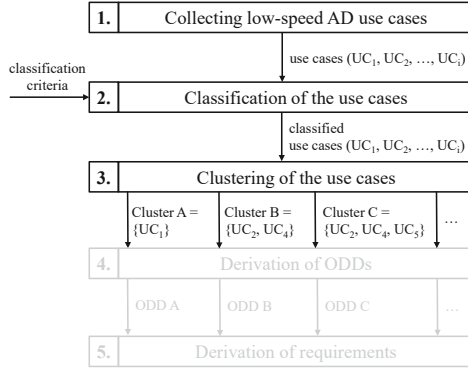


Figure 1: Methodology for the derivation of use case clusters for low-speed AD functions

2 Structure and Methodology

The applied methodology for the systematic derivation of the use case clusters for low-speed AD functions is shown in Figure 1. This methodology also defines the structure of this paper.

In the first step, described in Section 3, possible use cases and scenarios for low-speed AD functions are collected by considering two approaches – a *Stakeholder-Based Approach* and a *Velocity-Based Approach*. Subsequently, the found use cases are evaluated based on classification criteria. Therefore, in Section 4, classification criteria in two categories, covering the aspects of safety and technical capabilities, are first described and afterwards applied to the collected use cases. Based on this, the use cases are grouped into individual clusters that have common features regarding the evaluation of the individual criteria, which is described in Section 5. The identified clusters represent groups of use cases that each share common characteristics with respect to the evaluation criteria and therefore have the potential to be implemented through a generalized functional architecture. Subsequently, ODDs and requirements for specific applications can be derived from the single clusters. However, the latter mentioned two steps are not in the scope of this paper, but part of the development of the GLSF in the Project AUTotech.agil.

3 Collection of Low-Speed AD Use Cases

Two approaches are followed to collect possible use cases for low-speed AD functions as completely as possible. As a basis for the collection, the following constraints on low-speed AD functions are assumed to achieve the goal of a generalized low-speed AD function: The low-speed AD function represents a separate driving mode of the automated vehicle (AV), which is entered either from manual driving or from another AD mode. Thus, the low-speed AD function can be implemented both in SAE Level 3 and Level 4 systems. Consequently, there may or may not be occupants in the vehicle. Furthermore, only use cases are considered where the driving task is limited to transport on paved roads.

3.1 Stakeholder-Based Approach

In the first approach, the *Stakeholder-Based Approach*, different stakeholders are considered by investigating their possible needs and advantages regarding low-speed AD functions. In doing so, the following stakeholders are identified.

Firstly, a *vehicle owner* in the sense of a user of a private-owned vehicle, who is interested in the automation of challenging and annoying driving activities in the low-speed range. This results e.g. in *Pilot Services*, where the vehicle conditionally takes over the driving task (e.g. in a congestion). Furthermore, the use case group of *Automated Valet Services*, where the vehicle is automatically parked (AVP), charged (AVC), or washed (AVW), is of interest for the *vehicle owner*. Another stakeholder interested in this is a *municipality* (e.g. a city) that benefits from an efficient usage of the available parking space and reduced individual traffic for searching a parking spot.

A further stakeholder is summarized by the term *industry* and is aimed at all industrial applications, which automate activities that were previously carried out by humans. This results in the use case group of *Shuttle Services*, in which various entities are automatically transported from a defined starting point to a defined destination. Conceivable is the transport of people (*Group Shuttle*) or freight (*Cargo Shuttle*), for example on a large company site. Especially the *Group Shuttle* meets additionally the demands of a stakeholder *Commercial Operator*, which represents an entity managing a defined institution (e.g. airport, shopping centre, exhibition). Furthermore, the transport of the vehicle itself from the end of production line at the assembly plant to the transfer place is very resource-intensive and thus another use case for a low-speed AD function (*Self-Transportation Shuttle*), which is of special interest for the stakeholder *OEM*.

Another stakeholder is the *Automated Driving System (ADS)*, a central functional component in the AV that manages and ensures its correct and efficient functioning. In the case that the low-speed AD function represents a separate driving mode besides a main AD function, the following interests of the ADS in a low-speed AD function arise: Firstly, in the event of a degradation of the main AD system, the low-speed AD function can act as a fallback layer to bring the vehicle to a safe state. Furthermore, the low-speed AD function may have better characteristics in terms of maneuvering accuracy and energy consumption than the main AD system, so a change to the low-speed AD function may be reasonable in certain situations.

3.2 Velocity-Based Approach

In the second approach, the *Velocity-Based Approach*, possible use cases are collected considering all situations in road traffic where the vehicle's speed is below a threshold speed $v_{\max}$. In general, the particular value of $v_{\max}$ for low-speed AD functions is open. In this paper, a threshold speed $v_{\max}$ of 25 km/h is assumed, since a speed of 30 km/h already includes a high proportion of public road traffic, such as residential areas and main roads with reduced speed. In this case, the low-speed AD functions would not be distinguished from general AD functions for urban traffic.

The consideration of a threshold speed provides a variety of road traffic situations that can be classified into five groups, distinguished by the condition that induces the low-speed driving below the threshold speed. First, low-speed driving *induced by traffic*, e.g. congestion or right-of-way situations. Secondly, low-speed driving *induced by traffic*

regulations, which is dependant on the respective national road traffic regulations. In this work, the German Road Traffic Regulations (StVO) [19] are considered, resulting in situations such as driving through a traffic-calmed area (ger.: *Verkehrsberuhigter Bereich*). Thirdly, situations where low-speed driving is necessary due to a *scenery induced condition*, e.g. driving through a narrow street segment or manoeuvring into a parking lot. Fourthly, low-speed driving is necessary due to a *vehicle state induced condition*, such as a degradation of components of the AD system (see Stakeholder *ADS* in section 3.1). Fifth, the need of driving slowly due to a *vehicle load induced condition*, such as the transport of standing people. This situation can be linked directly to the use case *Group Shuttle* in section 3.1.

3.3 Summary and Functional Description

In summary, the found use cases can be divided into three overarching groups, as shown in Figure 2. Firstly, the **Pilot Services**, which refer to low-speed driving in public road traffic (e.g. congestion). This group is characterized by a frequent and condition-dependent transfer between „regular driving“ (other AD-mode or manual driving) and the low-speed use case. Furthermore, the possible environment of the use cases in this group can be a high number of road types with an unlimited variety of other road users. The second group are the **Automated Valet Services**, where the vehicle performs a certain service (parking, charging, washing) by AD in SAE Level 4 after the control was handed over to the Automated Driving System (ADS) at a defined drop-off zone and all occupants have left the vehicle. The drop-off zone can be either at the border of the area of the service (*internal drop-off*) or at an external place (e.g. airport terminal) (*external drop-off*). **Shuttle Services** represent the third group, which include the transfer of certain entities (people, freight, vehicle itself) from a defined starting point to a defined destination. In contrast to the Group *Pilot Services*, the groups *Automated Valet Services* and *Shuttle Services* are conducted by driving low speeds in delimited areas, which is

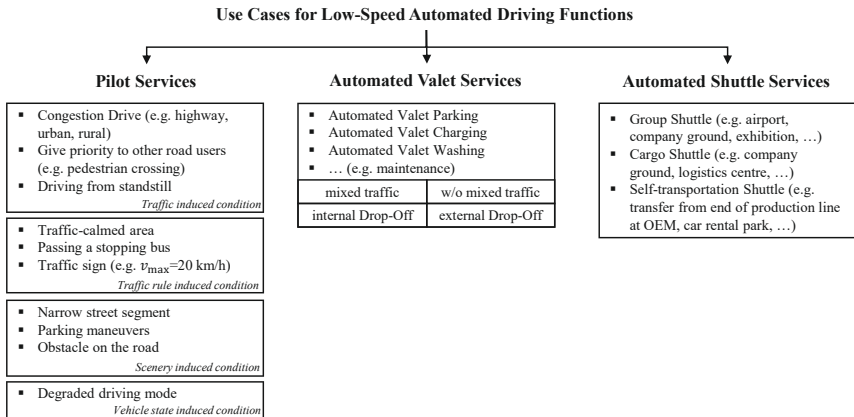


Figure 2: Overview of low-speed AD use cases

characterized by a well defined, e.g. position dependent, transfer into the low-speed use case. The environmental conditions and possible other road users are limited and can be predefined for the specific use case.

As a necessary prerequisite for the classification of the collected use cases (see Chapter 4), they are specified in two steps. First, all use cases whose functional characteristics may be included in another use case, are gathered together with the corresponding use case (e.g. *give priority to other road users* may be included in *urban congestion drive*). Second, for each use case a functional description is given. If necessary for the classification, constraints regarding the functional scope as well as the ODD of the corresponding use case are specified. The functional descriptions of the use cases considered in the following are given in Table 1. Each use case is identified with an ID (UC_X).

Use Case ID	Use Case	Functional description
UC _{1,1} (AVP) UC _{1,2} (AVC) UC _{1,3} (AVW)	Automated Valet Service (Parking/Charging/Washing) <i>w/o mixed traffic; internal drop-off</i>	The vehicle performs a certain service (parking, charging, washing, etc.) by Level 4 AD after control was handed over and all occupants have left the vehicle at a defined drop-off zone. Vehicle returns to defined pick-up zone after completing the service respectively upon user request. Pick-up/ drop-off zone can be <i>internal</i> (pick-up/ drop-off zone at border of area of the service) or <i>external</i> (pick-up/ drop-off zone outside of area of the service), which includes low-speed AD in public traffic to the area of the service. Inside the area of the service there either may only be AD vehicles (<i>w/o mixed traffic</i>) or people as well as manual driven cars (<i>with mixed traffic</i>) present.
UC _{2,1} (AVP) UC _{2,2} (AVC) UC _{2,3} (AVW)	Automated Valet Service (Parking/Charging/Washing) <i>with mixed traffic; internal drop-off</i>	
UC _{3,1} (AVP) UC _{3,2} (AVC) UC _{3,3} (AVW)	Automated Valet Service (Parking/Charging/Washing) <i>with mixed traffic; external drop-off</i>	
UC ₄	Group Shuttle	
UC ₅	Cargo Shuttle	Vehicle transfers certain entities (<i>Group Shuttle</i> : Sitting and standing people; <i>Cargo Shuttle</i> : Freight; <i>Self-transportation Shuttle</i> : Vehicle itself from end of production line at assembly plant to transfer place) from a defined starting point A to a defined ending point B by Level 4 AD.
UC ₆	Self-Transportation Shuttle	<i>Constraint</i> : Vehicle drives in delimited area (e.g. industrial site) between starting point and ending point; Only few instructed persons have access to transportation area (employees); No requirements to complex driving operations (e.g. parking maneuvers).
UC ₇	Parking Maneuver Assist	Function parks vehicle fully automated into a parking lot.
UC ₈	Narrow Segment Drive	Function takes over driving task for driving slowly through a narrow road segment with limited length (e.g. narrow underpass).
UC ₉	Traffic-Calmed Area Drive	Function takes over driving task in while driving through a traffic-calmed area (dt.: <i>Verkehrsberuhigter Bereich</i>).
UC ₁₀	Congestion Pilot (Urban)	Function takes over full driving task in urban area or on highway whenever a congestion is present and driven speed is below v_{max} .
UC ₁₁	Congestion Pilot (Highway)	<i>Constraints</i> : The AV will follow the vehicle ahead and execute no lane changes; Right of way is given by priority road or by traffic light (Congestion Pilot (urban)).
UC ₁₂	Speed Limit Drive	Function takes over full driving task whenever a speed limit below v_{max} is present (e.g. 20 km/h or 10 km/h).
UC ₁₃	Standstill Release	Function safeguards the start of a Level 4 AV in public road traffic until a threshold speed (e.g. 10 km/h) is reached by perceiving the environment with a near-field sensor module. <i>Constraints</i> : Function acts as a sense-only function.
UC ₁₄	Degraded Driving Mode	Function takes over driving task whenever a degradation of the regular architecture is detected and a minimal risk condition must be established <i>Constraints</i> : Function takes over below v_{max} (vehicle brakes until v_{max} is reached); Path of minimal risk maneuver is predefined (Function safeguards the predefined path); Environment is urban or highway with regular traffic and VRUs.

Table 1: Functional description of the low-speed AD use cases considered in this paper

4 Classification of the Use Cases

The goal of developing a generalized low-speed function implies that the requirements for the different use cases are completely fulfilled in one function. In addition to functional and technical requirements, the automotive industry places demands on safety, specifically defined by ISO 26262 (Functional Safety) [20] and ISO 21448 (SOTIF) [21]. Both set requirements depending on the risks associated with a function, e.g. expressed by an ASIL (automotive safety integrity level) in ISO 26262. Due to the diversity of the identified use

cases, for example regarding the complexity of the driving task or the accessibility of the environment for other road users, different risk and respective safety requirements and also different levels of required functional capabilities are expected. To be able to find intersections regarding the relevant requirements within the use cases, classification criteria are defined and assessed in the categories *safety requirements* and *technical capabilities*.

4.1 Safety Requirements

Under the aspect of safety, the risk associated with the execution of the use case with an AD function is assessed, to provide a tendency for the level of safety integrity requirements that will arise. This is important as it can be unfavorable to summarize a use case with low safety integrity requirements with a use case of high safety integrity requirements in one function. According to the state of the art, common risk parameters from ISO 26262 [20] are applied to evaluate the risk of a functional failure in the system. Since the specific architecture of the function is not yet known, however, the failure of safety-related components is not assessed explicitly. Rather, assuming any functional problem that causes a hazardous behaviour of the vehicle, the probability and severity of a collision is evaluated by using the following criteria:

- **Severity:** evaluating whether a collision is severe for a human life. This depends on whether the AV is occupied by humans as well as the type of other road including *unoccupied vehicles*, *occupied vehicles* and *vulnerable road users (VRU)*. Also, the maximum driving speed of the use case is considered, as e.g. driving with 10 km/h while parking leads to less harm than driving towards the end of a congestion with 25 km/h.
- **Exposure:** evaluating the presence of other road users, which depends on the traffic situation in the operative environment and can reach from *no other road users* over *occasional other road users* to *frequent traffic*.
- **Controllability:** evaluating the capabilities of the ego driver or other road users to react to the hazardous situation, e.g. with braking or with an evasive maneuver. This depends on the one hand on the possibility of the AV's occupants to manually control the vehicle, which is for example not possible in a group shuttle. On the other hand, available reaction times might differ depending on the use case.

Each criterion is evaluated on an ordinal scale including *low*, *medium* and *high*, leading to a corresponding final safety requirement level.¹

4.2 Technical Capabilities

The technical complexity for the realization of the function is evaluated by examining the required functional capabilities for the subfunctions of *sense*, *plan* and *act*. Since the functions are not yet fully technically defined, simplified assumptions are made. It thereby is assumed that the control and actuation stage is part of a common AV architecture, so that the low-speed function architecture is limited to the sensing and planning

¹It should be noted that these levels are not intended for deriving an ASIL rating, even if the same criteria are used for evaluation, but are only intended to establish comparability among the use cases.

level. Nevertheless, different use cases place different requirements on the precision of the executed trajectory, which is evaluated under the point *act*. Furthermore, it is assumed that a mission and/or route planning is also part of a central instance in the vehicle.

The classification of the required capabilities represents an estimate of the minimum technically necessary to accomplish the specific driving task, to keep the complexity level of the functions as low as possible. The following criteria are classified:

- **Sense capabilities:** classifying the required sensor field of views (FOV) in the direction (*front, sides, rear*) and distance (by means of *maximum speed*), required sensing quantities (*object size, position, relative speed*), object types or infrastructure elements to detect (e.g. *vehicles, pedestrians, cyclists, lane markings, traffic signs*) and level of perception (*only detection* or *classification*).
- **Plan capabilities:** classifying the stages or types of planning that are required to accomplish the AD task as well as the planning environment, as both can require different planning algorithms. The values include *behavior prediction, behavior planning, path and trajectory planning* and *parking maneuver planning*. The environment is distinguished between *structured, semi-structured* and *unstructured*.
- **Act capabilities:** classifying requirements in precisely following a planned path or goal position from *low* to *high*, as for example an AVC use case will place higher requirements on a final position precision than a Group Shuttle.

4.3 Exemplary Application

The classification of the use cases towards the aforementioned requirements is exemplarily shown for UC_{1.1} (*AVP w/o mixed traffic*), UC₄ (*Group Shuttle*), UC₉ (*Traffic-Calmed Area*) and UC₁₀ (*Congestion Pilot (urban)*) in Table 2. The classification of all use cases considered in this paper is given in the Appendix (Table 4 and Table 5).

UC_{1.1} is a use case with low safety requirements, which results from the low driving speeds of max. 12 km/h [9] as well as the fact that without mixed traffic no people are involved, neither in the vehicle nor in the environment. The perception requirements are therefore limited towards other vehicles and scenery elements. Assuming that the vehicle requires at least a continuous free space detection and can wait until a path is cleared, no classification or prediction of other objects is required.

UC₄ is evaluated with medium safety requirements, which is determined by the fact that standing persons can be inside the vehicle. Besides the increased safety level, the technical complexity is even lower than for UC_{1.1}, as no additional parking maneuver planner is required and the precision requirements are lower.

UC₉ is also a medium safety use case. One major difference from the previous use cases is that the ODD includes public roads. This means that the variety of other road users in the environment increases, yet a traffic-calmed area is only sparsely frequented and the driving speed is even limited to walking speed. The requirements for perception are increasing with regard to objects to be detected. An extreme example here is a ball that suddenly rolls onto the street and can be an indication of children playing. In addition, it is assumed that an object classification in this use case is needed as, in contrast to UC_{1.1} and UC₄, there are pedestrians in the environment to which the vehicle should keep a

larger safety distance than to scenery elements. Besides that, traffic-calmed areas still constitute an unstructured environment without lanes, traffic signs or right-of-way rules.

UC₁₀, the congestion pilot in urban environments, results in high safety requirements, which is due to the increased speed up to 25 km/h as well as frequent traffic in the surrounding area. It is assumed, for example, a main road with several successive traffic lights or road works that cause traffic jams. In addition to following a vehicle in front, the vehicle must therefore be able to perceive traffic control elements on the route, such as stopping at a traffic light or giving priority to pedestrians at a crosswalk. This results in additional perception tasks. Due to the increased maximum speed, the sensor field of view must be increased. Furthermore, a structured planning environment is present here in contrast to the other use cases.

		UC _{1,1} <i>AVP w/o mixed traffic, internal drop-off</i>	UC ₄ <i>Group Shuttle</i>	UC ₉ <i>Traffic-Calmed Area Drive</i>	UC ₁₀ <i>Congestion Pilot (urban)</i>
Safety Requirements	Severity	low (unoccupied vehicles)	medium (occupied vehicles, standing passengers)	medium (occupied vehicles, VRU, very low speed)	high (occupied vehicles, VRU, higher speed)
	Exposure	low (occasional other road users)	low (delimited area with occasional other road users)	low (occasional other road users)	high (frequent traffic in a congestion)
	Controllability	low (no human inside car)	low (no control elements in vehicle)	L3: medium (very low speed) L4: low (no intervention possible)	L3: low (small gaps → low time to react) L4: low (no intervention possible)
	Safety level	low	medium	medium	high
Technical Requirements	Sensor FOV	front, sides, rear (very small blind zone req.)	front, sides (no lane changes, no backward driving)	front, sides (no lane changes, no backward driving)	front, sides (no lane changes, no backward driving)
	Sensing distance → max. speed	12 km/h	12 km/h	7-8 km/h "walking speed"	25 km/h
	Objects to detect	scenery obstacles, vehicles in path (ahead, oncoming or crossing), empty parking space	scenery obstacles, vehicles in path, pedestrians (employees)	scenery obstacles, vehicles, pedestrians, cyclists & MVs ahead or oncoming or crossing, play tools like ball	scenery obstacles, vehicles, pedestrians, cyclists & MVs ahead or oncoming or crossing, lanes, traffic signs
	Sensing quantities	ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects	ego position, size and relative position and speed of relevant objects, lanes, traffic sign
	Object classification	no	no	yes (e.g. to keep more safety buffer towards pedestrians)	yes (e.g. giving priority to pedestrians)
	Prediction	no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)
	Planning task	path & trajectory, parking	path & trajectory	path & trajectory	path & trajectory
	Planning environment	unstructured / semi structured	semi-structured	unstructured / semi structured	structured
	Precision req.	high req.	low req.	low req.	low req.

Table 2: Application of the classification scheme to selected low-speed AD use cases

5 Clustering of the Use Cases

The clustering process searches for overlaps in the safety and technical requirements, classified according to the procedure described above. This can result in smaller use case clusters with a low requirements threshold or high but very specific requirements, as well as larger use case clusters that cover simple and complex use cases. This reveals different possibilities of a GLSF, depending on the maximum safety or technical complexity level that is chosen.

The starting point for the clustering are all use cases with a low safety level and matching, low technical requirements, which then form the first "elementary" clusters.

Afterwards, the thresholds in the dimensions of safety and technical capabilities are incrementally increased and partly overlaps to existing clusters are searched, expanding previous clusters. Thereby, further "elementary" clusters with higher requirement levels that have no overlaps to other previous clusters are identified as well.

5.1 Exemplary Application

To demonstrate the methodology outlined above, it is performed exemplarily and a resulting set of related use case clusters is described. This yields the use case clusters A.0, A.1, A.2 and A.3, whose properties are shown in Table 3 for a better understanding of the following.

ID	Use Case Cluster	Use Cases	Properties
A.0	Valet Services w/o mixed traffic	AVP/AVC/AVW w/o mixed traffic	Safety requirements: low Technical properties: max. 12 km/h; 360° perception required; delimited un-/semi-structured environment; parking maneuvers required
A.1	Valet and Shuttle Services w/o human transport	A.0; Cargo Shuttle; Self-Transportation Shuttle	Safety requirements: low Technical properties: see A.0; detection of humans required (rare presence of humans)
A.2	Valet and Shuttle Services with human transport	A.1; Group-Shuttle; AVP/AVC/AVW with mixed traffic; Parking Maneuver Assist	Safety requirements: medium Technical properties: see A.1; increased requirements to object detection capabilities (e.g. object classification) and response time (frequent presence of humans in direct surrounding possible)
A.3	Valet and Shuttle Services with low-speed public traffic extension	A.2; Traffic-Calmed Area Drive; Narrow Segment Drive; Standstill Release	Safety requirements: medium Technical properties: see A.2; extension to public traffic in semi-structured environment (Traffic-Calmed Area)

Table 3: Extract of derived low-speed AD use case clusters

For the first cluster, use cases with low safety requirements and similar low technical requirements are aggregated from all classified low-speed AD use cases (see Tables 4 and 5, Appendix). Thereby, the "elementary" cluster A.0 (*Valet Services w/o mixed traffic*) is obtained, which includes the three valet services AVP, AVC, and AVW under the boundary condition *no mixed traffic* (UC_{1.1}, UC_{1.2}, UC_{1.3}). Due to the lack of human presence in the ODD of the use cases, all three use cases have low safety requirements. Furthermore, the technical requirements are almost equivalent due to the functional similarity of the three use cases. Therefore a cluster is formed. The technical characteristics of cluster A.0 are shown in Table 3. By increasing the technical requirements while keeping the safety requirements constant, cluster A.1 (*Valet and Shuttle Services w/o human transport*) is formed. In addition to the use cases from cluster A.0, this includes the use cases *Cargo Shuttle* and *Self-Transportation Shuttle* (UC₅, UC₆). In this case, the increase in technical requirements is due to the need for a detection of humans in the environment. The remaining technical requirements for the added use cases are already covered in cluster A.0. After the increase of the technical requirements, an increase of the safety requirements to the level *medium* takes place with the technical requirements remaining as constant as possible, resulting in cluster A.2 (*Valet and Shuttle Services with human transport*). The increased safety requirements here result from the transportation of people in the AD vehicle in the use case Group Shuttle (UC₄), which otherwise has no further technical requirements compared to the previous cluster A.1. The increase in safety requirements as well as the already existing technical requirements further enable the addition of the use case group of valet services with mixed traffic (UC_{2.1}, UC_{2.2}, UC_{2.3}), which in turn results in slightly increased technical requirements (object detection and

classification). Due to the necessary safety requirements and technical requirements, the addition of the use case *Parking Maneuver Assist* (UC₇) is also possible. In the last step, the technical requirements are further increased while the safety requirements remain the same. This results in cluster A.3 (*Valet and Shuttle Service with low-speed public traffic extension*), in which the use cases *Traffic-Calmed Area Drive* (UC₉), *Narrow Segment Drive* (UC₈) and *Standstill Release* (UC₁₃) are added. The technical requirements increase due to the extension of the use cases to public road applications, which has an impact on the necessary perception and planning capabilities.

In summary, cluster group A (A.0 to A.3) covers all low-speed AD use cases in delimited areas that are characterized by the same maximum speed of 12 km/h and a similar planning task. Cluster A.3 also extends to use cases on public road traffic with the same maximum speed and still limited technical requirements (e.g. no lane detection necessary).

The complete overview of the defined use case clusters is shown in the appendix (Table 6). In addition to the cluster group A, four further cluster groups were defined. Cluster groups B and C each form subsets of the already described group A. To be highlighted is cluster group D (D.0 to D.2), which is characterized by high safety requirements and technical properties that diverge from cluster group A (e.g. increased maximum speed, behavior planning, traffic sign and lane detection) resulting from applications in public road traffic. The mentioned differences entail that the use cases from cluster group D cannot be combined directly with the use cases from cluster group A. A special case is the use case of Valet Services with external drop-off (UC_{3.1}, UC_{3.2}, UC_{3.3}), whose technical requirements include all technical requirements of the remaining use cases due to the ODD, which includes both public road traffic and delimited areas. This results in use case cluster E (*Holistic low-speed function*), which includes all use cases defined in this paper. If all vehicle concepts of the disruptive modular platform architecture of the UNICAR.agil project (private vehicle, taxi, group shuttle, cargo shuttle) [18] were to be considered in one function, cluster E would be the starting point of development.

5.2 General Findings

In addition to the clusters itself, further conclusions can be drawn from the classification and clustering process: It is observable that still different, distinct speed levels exist within the low-speed domain (up to 12 km/h and up to 25 km/h), leading to varying requirements in e.g. safety or sensor FOV. It is thereby found that distinct elemental clusters of use cases can be built, while it is observed that lower safety requirements often correlate with lower technical demands. Furthermore, the variety of planning environments (structured to unstructured) and planning tasks (parking, free drive, follow-up drive) encountered in different use cases underscores the need for a GLSF that can handle a wide range of tasks and environments when choosing a large cluster such as Cluster E. It can also be seen that there is no continuous hierarchy within the use cases with regard to requirements. In the case of larger use case clusters, the highest occurring requirements per category are not only contributed from one, but from different use cases. The exception here are the previously mentioned Valet Services with external drop-off. Finally, the aggregation of use cases into clusters not only aids in organizing requirements but also presents opportunities for expanding the functional scope of individual use cases. For instance, combining Valet Services with Shuttle Services can enhance the functionality of the latter, leveraging existing parking capabilities.

6 Conclusion and Outlook

In this paper, a methodology is presented to systematically derive potential use cases for low-speed AD functions and to cluster them with respect to applicable requirements, so that use case clusters for a generalized low-speed AD function can be inferred according to defined technical constraints.

The paper contributes to the future development of low-speed AD functions in two dimensions: First, a structured and holistic overview of potential application areas for AD at low speeds is provided, enabling the development of future research areas. Second, future development of generalized low-speed AD functions is facilitated because the ODD underlying a function to be developed, which forms the basis for the further definition of the function's requirements and architecture, can be defined in a simplified manner using the use case clusters derived in this paper. By defining the clusters according to the dimensions of *Safety* and *Technical Capabilities*, it is possible to align the selection of clusters with a developer's own constraints for the desired function development.

The proposed use cases apply only to a speed range up to 25 km/h. Accordingly, an extension of the speed range may lead to further use cases, which will allow to enlarge or form new clusters using the presented methodology.

It should be noted that the results obtained are influenced by certain specifications regarding the boundary conditions of the use cases. For example, expanding the maximum speed of the Group Shuttle Services can lead to different cluster results. Also, assumptions towards required technical capabilities of the functions were made. Therefore, ongoing implementation and validation is essential to ensure that the estimated technical capabilities, which served as minimum requirements, are indeed sufficient to obtain reasonable behavior of the functions in real world applications.

Based on the rough requirements definition presented here, high-level architectures for a GLSF can already be derived for selected clusters. The next step for the development of a GLSF is the concretization of the ODD that results from the use case clusters. Based on this, requirements can be concretized within the categories shown by deriving the worst-case framework conditions of the ODD. These are then used, for example, to derive a suitable sensor module for the short range as well as suitable planner metrics and algorithms.

Acknowledgments

This research is accomplished within the project "AUTOftech.agil" (FKZ 01IS22088S). We acknowledge the financial support for the project by the Federal Ministry of Education and Research of Germany (BMBF).

References

- [1] SAE International, "SAE International Standard J3016: Taxonomy and Definitions for Terms related to On-Road Motor Vehicle Automated Driving Systems," 2014.

- [2] H. Winner, “ADAS, Quo Vadis?,” in *Handbook of Driver Assistance Systems* (H. Winner, S. Hakuli, F. Lotz, and C. Singer, eds.), Springer International Publishing, 2016.
- [3] M. Bolle, S. Knoop, F. Niewels, and T. Schamm, “Early level 4/5 automation by restriction of the use-case,” in *17. Internationales Stuttgarter Symposium* (M. Barge, H.-C. Reuss, and J. Wiedemann, eds.), Proceedings, pp. 531–545, Wiesbaden: Springer Fachmedien Wiesbaden, 2017.
- [4] Kraftfahrt-Bundesamt, “KBA erteilt erste Genehmigung für fahrerlos einparkendes Fahrzeug,” 30.11.2022. https://www.kba.de/DE/Presse/Pressemitteilungen/Allgemein/2022/pm45_2022_AVP_erste_Genehmigung.html, visited 10.05.2023.
- [5] Mercedes-Benz Group AG, “Easy Tech: Conditionally automated driving with the DRIVE PILOT,” 2021. <https://group.mercedes-benz.com/company/magazine/technology-innovation/easy-tech-drive-pilot.html>, visited 24.05.2023.
- [6] H. Banzhaf, D. Nienhuser, S. Knoop, and J. M. Zollner, “The future of parking: A survey on automated valet parking with an outlook on high density parking,” in *2017 IEEE Intelligent Vehicles Symposium (IV)*, pp. 1827–1834, IEEE, 2017.
- [7] M. Müller, “Connected Parking,” *Bautechnik*, vol. 94, no. 5, pp. 313–318, 2017.
- [8] Kraftfahrt-Bundesamt, “Technischer Anforderungskatalog für die autonome Fahrfunktion „Automated Valet Parking (AVP)“,“ 2022. https://www.kba.de/DE/Themen/Typgenehmigung/Zum_Herunterladen/autonomes_automatisiertes_Fahren/dl_anforderungskatalog_AVP.pdf, visited 24.05.2023.
- [9] ISO, “ISO/FDIS 23374-1: Intelligent transport systems — Automated valet parking systems (AVPS) — Part 1: System framework, requirements for automated driving and for communications interface,” Under Development.
- [10] ISO, “ISO/AWI 12768: Intelligent transport systems — Automated Valet Driving Systems (AVDS),” Under Development.
- [11] U. Schwesinger, M. Burki, J. Timpner, S. Rottmann, L. Wolf, L. M. Paz, H. Grimmett, I. Posner, P. Newman, C. Hane, L. Heng, G. H. Lee, T. Sattler, M. Pollefeys, M. Allodi, F. Valenti, K. Mimura, B. Goebelsmann, W. Derendarz, P. Muhlfellner, S. Wonneberger, R. Waldmann, S. Gryczyk, C. Last, S. Bruning, S. Horstmann, M. Bartholomäus, C. Brummer, M. Stellmacher, F. Pucks, M. Nicklas, and R. Siegwart, “Automated valet parking and charging for e-mobility,” in *2016 IEEE Intelligent Vehicles Symposium (IV)*, pp. 157–164, IEEE, 6/19/2016 - 6/22/2016.
- [12] Project Autoples, “Autoples – Automatisiertes Parken & Laden von Elektrofahrzeug-Systemen,” 2015.
- [13] Embotech, “Full-scale solution: Automated Parking, Driving, Charging & Washing as one service sequence,” 2023. https://www.youtube.com/watch?v=y_ax1q04Udo&ab_channel=embotech, visited 24.05.2023.

- [14] J. Cregger, M. Dawes, S. Fischer, C. Lowenthal, E. Machek, and D. Perlman, “Low-Speed Automated Shuttles: State of the Practice Final Report,” 2018.
- [15] BMW Group, “Pilot project: Cars manoeuvre in production without drivers,” 2022. <https://www.press.bmwgroup.com/global/article/attachment/T0402335EN/564387>, visited 24.08.2023.
- [16] ISO, “ISO 22737:2021: Intelligent transport systems — Low-speed automated driving (LSAD) systems for predefined routes — Performance requirements, system requirements and performance test procedures,” 2021.
- [17] AUTOtech.agil, “AUTOtech.agil: Architektur und Technologien zur Orchestrierung automobiltechnischer Agilität,” 2022. https://www.fzd.tu-darmstadt.de/forschung/research_projects_fzd/rp_autotechagil/index.en.jsp, visited 15.05.2023.
- [18] T. Wooten, L. Eckstein, S. Kowalewski, D. Moormann, M. Maurer, R. Ernst, H. Wimmer, S. Katzenbeisser, M. Becker, C. Stiller, K. Furmans, K. Bengler, M. Lienkamp, H.-C. Reuss, K. Dietmayer, H. Latagahn, N. Siepenkötter, M. Elbs, E. v. Hinüber, M. Dupuis, and C. Hecker, “UNICARagil - Disruptive Modular Architectures for Agile, Automated Vehicle Concepts,” in *27. Aachener Kolloquium Fahrzeug- und Motorentechnik*, pp. 663–694, 2018.
- [19] Bundesministerium der Justiz und für Verbraucherschutz, “Straßenverkehrs-Ordnung (StVO),” 6. März 2013.
- [20] ISO, “ISO 26262-1: Road vehicles — Functional safety — Part 1: Vocabulary,” 2018.
- [21] ISO, “ISO 21448: Road vehicles — Safety of the intended functionality,” 2022.

Appendix

Safety Requirements		UC _{1,1} , UC _{1,2} , UC _{1,3} <i>AVP, AVC, AVW w/o mixed traffic, internal drop-off</i>	UC _{2,1} , UC _{2,2} , UC _{2,3} <i>AVP, AVC, AVW with mixed traffic, internal drop-off</i>	UC _{3,1} , UC _{3,2} , UC _{3,3} <i>AVP, AVC, AVW with mixed traffic, external drop-off</i>	UC ₄ <i>Group Shuttle</i>	UC ₅ , UC ₆ <i>Cargo Shuttle, Self-Transportation Shuttle</i>	UC ₇ <i>Parking Maneuver Assist</i>	UC ₈ <i>Narrow Segment Drive</i>
Severity		low (unoccupied vehicles)	medium (other road users, VRU, very low speed)	high (other road users, VRU)	medium (occupied vehicle, standing passengers)	low (low speed and no humans inside vehicle)	low (very low speed)	medium (other road users, VRU, very low speed)
Exposure		low (occasional other road users)	low (occasional other road users)	high (frequent traffic like crossing pedestrians or other vehicles ahead)	low (delimited area with occasional other road users)	low (delimited area with occasional other road users)	low (few other road users during parking)	Low (few other road users during maneuver)
Controllability		low (no human inside car)	low/medium (no human inside car, but reaction of pedestrians due to low speed possible)	low (no human inside car)	low (no control elements in vehicle)	low/medium (no human inside car; humans with system knowledge around vehicle)	L3/L4: medium/low (hum. interv., possible w. small react. time / no hum. intervention possible)	L3/L4: high/low (human intervention possible / no human intervention possible)
Safety level		low	medium	high	medium	low	low	medium
Sensor FOV		front, sides, rear (very small blind zone req.)	front, sides, rear (very small blind zone req.)	front, sides, rear (very small blind zone req.)	front, sides (no lane changes, no backward driving)	front, sides (no lane changes, no backward driving)	front, sides, rear (very small blind zone req.)	front, sides, rear (very small blind zone req.)
Sensing distance (max. speed)		12 km/h	12 km/h	25 km/h	12 km/h	12 km/h	5 km/h	12 km/h
Objects to detect		scenery obstacles, vehicles in path ahead, oncoming or crossing, empty parking space	scenery obstacles (e.g. shopping trolley), vehicles in path ahead, oncoming or crossing, empty parking space, pedestrians	scenery obstacles (e.g. shopping trolley), vehicles in path ahead, oncoming or crossing, empty parking space, lanes, vehicle ahead, pedestrians, cyclists, crossing paths with priority, traffic signs	scenery obstacles, vehicles in path, pedestrians (employees)	scenery obstacles, vehicles in path, pedestrians (employees)	scenery obstacles, other vehicles or persons next to/in parking lot	all kinds of objects very close to vehicle, esp. in path
Sensing quantities		ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects, relative speed to vehicle ahead, sign recognition	ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects	ego position, size and relative position of relevant objects
Object classifc.		no	yes	yes	no	no	no	no
Prediction		no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)	yes (e.g. lane changes should be possible)	no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)	no	no
Planning task		path & trajectory, parking	path & trajectory, parking	path & trajectory, parking (evtl. behavior planning)	path & trajectory	path & trajectory	parking planner	path & trajectory
Planning env.		unstructured / semi structured	unstructured / semi structured	unstructured / semi structured & structured environment	semi-structured	semi-structured	semi-structured	(semi)-structured
Precision req.		high req. (AVC: 20 cm (conductive), 2 cm (inductive))	high req. (AVC: 20 cm (conductive), 2 cm (inductive))	high req. (AVC: 20 cm (conductive), 2 cm (inductive))	low req.	low req.	high req.	high req.

Table 4: Application of the classification scheme to the low-speed AD use cases (I/II)

Safety Requirements						UC ₁₄ <i>Degraded Driving Mode</i>
Severity	UC ₉ <i>Traffic-Calm Area Drive</i>	UC ₁₀ <i>Congestion Pilot (Urban)</i>	UC ₁₁ <i>Congestion Pilot (Highway)</i>	UC ₁₂ <i>Speed Limit Drive</i>	UC ₁₃ <i>Standstill Release</i>	(other road users, VRU) high
Exposure	medium (occupied vehicles, VRU, very low speed)	high (occupied vehicles, VRU, higher speed)	medium (no VRUs, higher speed)	high (other road users, VRU)	medium (other road users, VRU, very low speed)	high (other road users, VRU)
Controllability	low (occasional other road users)	high (frequent traffic in a congestion)	high (frequent traffic in a congestion)	high (frequent traffic like pedestrians or other vehicles ahead)	high (frequent traffic like crossing pedestrians or other vehicles ahead)	high (frequent traffic like pedestrians or other vehicles ahead)
Safety level	L3: medium (very low speed) L4: low (no intervention possible)	L3: low (small gaps → low time to react) L4: low (no intervention possible)	L3: low (small gaps → low time to react) L4: low (no intervention possible)	L3/L4: high/low (human intervention possible/ no human intervention possible)	low (no human intervention possible)	low (no human intervention possible)
Sensor FOV	medium	high	medium	high	medium	high
Sensing distance (max. speed)	front, sides (no lane changes, no backward driving)	front, sides (no lane changes, no backward driving)	front, sides (no lane changes, no backward driving)	front, sides, rear (lane changes possible)	front & sides (no backward driving)	front & sides (only safeguarding of predefined path)
Objects to detect	7-8 km/h "walking speed"	25 km/h	25 km/h	20 km/h (first traffic sign below 30 km/h)	0-10 km/h	25 km/h
Sensing quantities	scenery obstacles, vehicles, pedestrians, cyclists & MVs ahead or oncoming or crossing, play tools like ball	scenery obstacles, vehicles, pedestrians, cyclists & MVs ahead or oncoming or crossing, lanes, traffic signs	lanes, vehicles ahead	lanes, vehicle ahead, pedestrians, cyclists, motor vehicles, crossing paths with priority, traffic signs, lane detection	all kinds of objects very close to vehicle, esp. in path, → very small blind zone	all kinds of collision objects and obstacles in path
Object classific.	ego position, size and relative position of relevant objects	ego position, size and relative position and speed of relevant objects, lanes, traffic sign	ego position, distance & relative speed to vehicle ahead	ego position, distance & relative speed to vehicle ahead, distance and size of relevant objects, sign recognition	ego position, size and relative position of relevant objects, sign recognition	ego position, size and relative position of relevant objects in path
Prediction	yes (e.g. keep more safety buffer towards pedestrians)	yes (e.g. giving priority to pedestrians)	no	yes	no	no
Planning task	no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)	no (vehicle will wait until path is cleared)	yes (complex maneuvers like lane changes possible)	sense-only function	no (vehicle will wait until path is cleared)
Planning env.	path & trajectory	path & trajectory	path & trajectory	behaviour planning, path & trajectory		trajectory
Precision req.	unstructured / semi structured	structured	structured	structured	sense-only function	path is preplanned
	low req.	low req.	low req.	low req.		low req.

Table 5: Application of the classification scheme to the low-speed AD use cases (II/II)

ID	Use Case Cluster	Use Cases	Properties	Stakeholders
A.0	Valet Services w/o mixed traffic	AVP/AVC/AVW w/o mixed traffic	Safety requirements: low Technical properties: max. 12 km/h; 360° perception required; delimited un-/semi-structured environment; parking maneuvers required	Vehicle owner
A.1	Valet and Shuttle Services w/o human transport	A.0; Cargo Shuttle; Self-Transportation Shuttle	Safety requirements: low Technical properties: see A.0; detection of humans required (rare presence of humans)	Vehicle owner; Industry; OEM
A.2	Valet and Shuttle Services with human transport	A.1; Group-Shuttle; AVP/AVC/AVW with mixed traffic; Parking Maneuver Assist	Safety requirements: medium Technical properties: see A.1; increased requirements to object detection capabilities (e.g. object classification) and response time (frequent presence of humans in direct surrounding possible)	Vehicle owner; Industry; OEM; Comm. Operators
A.3	Valet and Shuttle Services with low-speed public traffic extension	A.2; Traffic-Calmed Area Drive; Narrow Segment Drive; Standstill Release	Safety requirements: medium Technical properties: see A.2; extension to public traffic in semi-structured environment (<i>Traffic-Calmed Area</i>)	Vehicle owner; Industry; OEM; Comm. Operators; ADS
B.0	Very/low speed public traffic	Traffic-Calmed Area Drive; Narrow Segment Drive; Standstill Release	Safety requirements: medium Technical properties: max. 12 km/h; no 360° perception required; semi-structured environment; high requirements to object detection capabilities and response time (frequent presence of humans in direct surrounding possible); no parking maneuvers required	Vehicle owner; ADS
C.0	Shuttle Service w/o human transport	Cargo Shuttle; Self-Transportation Shuttle	Safety requirements: low Technical properties: max. 12 km/h; no 360° perception required; delimited un-/semi-structured environment; detection of humans required (rare presence of humans); no parking maneuvers required	Industry; OEM
C.1	Shuttle Service with human transport	C.0; Group-Shuttle	Safety requirements: medium Technical properties: see C.0	Industry; OEM; Further Operator
C.2	Shuttle Service with low-speed public traffic extension	C.1; B.0	Safety requirements: medium Technical properties: see C.0; increased requirements to object detection capabilities and response time (frequent presence of humans in direct surrounding possible)	Industry; OEM; Comm. Operators; ADS
D.0	Highway pilot	Congestion Pilot (Highway); Degrade Driving Mode	Safety requirements: high Technical properties: max. 25 km/h; no 360° perception required; lane detection required; measurement of relative speed required; structured environment; simple planning task	Vehicle owner; ADS
D.1	Urban und highway pilot	D.0; Congestion Pilot (Urban)	Safety requirements: high Technical properties: see D.0; increased perception requirements (traffic sign detection, detection of further objects, object classification); higher planning requirements (giving priority to other road users)	Vehicle owner; ADS
D.2	Urban and highway pilot extended	D.1; C.1; B.0; Speed Limit Drive	Safety requirements: high Technical properties: see D.1; 360° perception required; higher planning requirements (behavior planning, behavior prediction)	Vehicle owner; Industry; OEM; Comm. Operators; ADS
E	Holistic low-speed function	All Use Cases	Safety requirements: high Technical properties: Function covers technical requirements of all mentioned use-cases of low-speed AD functions due to including use case AVP/AVC/AVW with external drop off	All mentioned Stakeholders

Table 6: Overview of low-speed AD use case clusters

Derivation of quantitative risk acceptance criteria for automated driving systems

Jan Erik Stellet, Bernd Müller, Susanne Ebel*

Abstract: Defining risk acceptance criteria for automated driving systems is an essential step for a successful release and avoidance of field incidents. Despite several regulatory provisions and normative frameworks, there is not yet a common understanding and approach, particularly concerning quantitative criteria. This work firstly gives a structured analysis of requirements and approaches. The second contribution is the proposal of a new approach targeting effectively no fleet incidents (ENFLI).

Keywords: Automated Driving, Safety, SOTIF, risk acceptance

1 Motivation

The safety of a product, i.e., not causing harm, is one (although not the only) crucial property for its persistent success on the market and avoidance of legal risks for the manufacturer. However, since perfect safety is typically not achievable, the question what defines acceptable safety risks, i.e., risk acceptance criteria (RAC), arises.

For the safety assurance of automated driving systems (ADS), the definition of defensible (quantitative) RAC is a highly relevant topic. Technical regulations, e.g., UNECE R157 [1] or EU 2022/1426 [2] include high-level provisions yet do not prescribe specific criteria. Frameworks and guidance are provided in industry standards such as ISO 21448 [5] or the upcoming ISO TS 5083. Several publications, e.g., [6–8] address gaps and questions regarding the practical application of these frameworks. However, despite these efforts, a common understanding and accepted approach has not been reached yet.

To advance the discussion, this paper provides context on the applicability of quantitative RAC (Sec. 2). It then summarises regulatory and normative requirements on RAC (Sec. 3) as well as commonly known approaches for the derivation of quantitative RAC (Sec. 4). To address an observed inconsistency between these approaches and the safety level achieved by mature automotive systems, a new approach is proposed, exemplified, and discussed (Sec. 5).

2 Context on quantitative risk acceptance criteria

While safety and risk acceptance are a very broad topic, this paper will focus on quantitative RAC. They typically are to be interpreted within some context and refer to some but not all safety relevant properties of a system. It is crucial to keep the limitations on the

*The authors are with Robert Bosch GmbH, Cross-Domain Computing Solutions, Stuttgart, Germany.

scope and potential purpose of quantitative RAC in mind when evaluating approaches for their applicability in the ADS domain.

In the understanding adopted in this work, quantitative RAC address only hazards at the vehicle level.¹ From these RAC, further quantities can be derived that can also relate to lower levels of an ADS system architecture. According to ISO 21448 [5], verification and validation (V&V) targets “provide evidence that the (risk) acceptance criteria are met”. For an extensive analysis of the use of RAC and other quantities we refer to [7].

There are different purposes of quantitative RAC that can be considered but not all purposes can be addressed equally well as is shown in Fig. 1. This underlines that quantitative RAC are only a subset of a more holistic set of acceptance criteria.

A quantitative risk acceptance criterion refers to the number of some (countable) critical event. While there are multiple options, some, like the number of accidents with fatalities, are lagging measures that can be only applied after a product has been introduced to the market. They are not directly useful for product engineering. Leading measures, like the number of safety goal violations (ISO 26262) or occurrences of hazardous behaviour (ISO 21448), are a more appropriate choice.

Furthermore, a quantitative criterion addresses the acceptability of consequences of faults (including functional insufficiencies and their consequences, cf. ISO 21448 [5]), failures or malfunctions. However, for some faults, quantitative RAC are not applicable at all, e.g., classical software bugs, misuse scenarios or security related attacks.

Finally, every quantitative criterion is only meaningful with reference to the measurement principle with which it is evaluated. Generalised conclusions must be handled with care. Usefulness for safety considerations essentially depends on, e.g., whether all circumstances that allow an observation of the failure phenomenon are sufficiently addressed by the measurement principle.

3 Requirements from regulations and standards

There are ADS type approval regulations which contain provisions on (quantitative) RAC, namely UNECE R157 [1], EU 2022/1426 [2] and the German AFGBV² [3]. Furthermore, the standard ISO 21448 [5] addresses the safety of the intended functionality (SOTIF) for E/E systems where proper situational awareness is essential to safety derived from complex sensors and processing algorithms.

Tab. 1 highlights similarities and differences in these texts with respect to RAC. Note that this is based on the authors’ interpretation of partly differing terminologies.

¹Note that, while this understanding is in line with ISO 21448 [5, clause 6.1], a closer look at Annex C 2.1 of the standard shows that there can be more than one RAC on the vehicle level. In fact, the example mentions RAC with three different scopes namely (1) an “original acceptance criterion” which is used as an aggregate figure over all accident/incidents, (2) an acceptance criterion for individual accidents/incidents, and (3) following a decomposition of the former to hazardous behaviours, an “acceptance criterion of this behaviour”.

²*Autonome-Fahrzeuge-Genehmigungs-und-Betriebs-Verordnung*, engl.: German implementing ordinance for automated and autonomous vehicles

Purpose	Addressable with quantitative RAC
Compliance to type approval regulations & safety standards	▪ Required by SAE L3 and L4 type approval regulations. ▪ Quantitative RAC <i>can be used</i> according to ISO 21448. ▪ Other safety standards rely on mostly qualitative measures. → Depending on applicable regulation or standard ■ ■ □
Social, market or legal acceptance	Public and legal reaction to statistically similar field incidents varies strongly. → Limited usefulness of of quantitative RAC □ □ □
Demonstration of safety / safety case	Typical safety case consists of many qualitative and quantitative arguments & evidences. → A contribution but never the whole demonstration of safety ■ □ □
Derivation of architecture and design targets	Safety is not a directly measurable property, but quantitative targets can help in evaluating architectures & designs. → Very useful application of quantitative (vehicle-level) RAC ■ ■ ■
Derivation of V&V targets	V&V targets can be better argued if derived from RAC. → Very useful application of quantitative RAC ■ ■ ■

□ □ □ Not addressable ■ ■ ■ Well addressable

Figure 1: Potential purposes and how well they can be addressed by quantitative RAC.

Table 1: Summary of provisions on quantitative RAC.

	Type approval regulations			Standard
	UNECE R157	EU 2022/1426	AFGBV	ISO 21448
Types of risks in scope:	Functional safety and SOTIF-related			SOTIF-related
Use of qualitative or quantitative RAC:	Qualitative and quantitative			Qualitative or quantitative or both
Principle for the derivation of quantitative RAC	Comparison to human driver:			Examples given
	“Unreasonable risk means the overall level of risk ... compared to a competently and carefully driven manual vehicle” [1, Annex 4, 2.16]	“Unreasonable risk means the overall level of risk ... compared to a manually driven vehicle in comparable transportation services and situations” [2, Article 2, 28.]	“Maß an Sicherheit ... höher als das Maß an Sicherheit bei Fahrzeugen, die von Personen geführt werden” [3, Anlage 1, 10.]	include selection or combination of the principles outlined in Sec. 4. A rationale is required. [5, clause 6.5]

Table 2: Comparison of commonly referred to principles for the derivation of RAC.

Principle	Summary	Reference of risk	Acceptance criteria
MEM	Acceptable risk is calculated based on the lowest rate of mortality for human individuals in the general population.	The lowest rate of mortality for human individuals	The individual risk (fatalities per person and time) caused by the system is lower than the tolerable risk derived from MEM.
ALARP	Risks are separated in three bands: 1) intolerable, 2) ALARP, 3) broadly acceptable. Individual risks in band 2) are reduced to a level considered “reasonably practicable” by weighing the risk against the effort needed to further reduce it.	The change in collective risk associated with each option/ safety measure.	If the costs of a measure are judged to be disproportionate to the safety benefits, then the measure is judged not to be necessary to further reduce the risk. Several factors need to be considered in this judgement, e.g., state of the art.
GAMAB /GAME	Comparison of two systems: the new system must be globally as safe as or safer than the existing one.	Reference system – could be the human driver if no comparable technical system exists.	The new system is less or equally risky compared to the existing system.
PRB	Allows counterbalancing of residual risks against safety benefit.		

4 Survey of existing approaches

Several approaches for the derivation of RAC have been developed in different application areas. In the railway domain the definition of risk acceptance criteria is described in the CENELEC safety standard EN 50126. There, the principles MEM (Minimum endogenous mortality), ALARP (As low as reasonably practicable) and GAMAB/GAME (Globalment au moins bon / équivalent) are described as methods to define risk acceptance criteria. In addition, the Positive risk balance (PRB) argumentation is mentioned in the code of ethics by the German national ethics committee on automated and connected driving. Tab. 2 provides a basic comparison of these principles.

In the MEM principle, the reference of risk is independent of the technology to be developed. However, it stands to reason to use a reference of risk that is more specific for the automated driving domain (e.g., human traffic fatal accident rates) and therefore, the acceptability of the generic MEM value is questionable. The same applies to the ALARP principle, where quantitative values known from the literature for separation between the three bands of risks are unspecific for ADS.

In order to use GAME/GAMAB for the introduction of new ADS it seems obvious to

refer to human driver accident statistics as easily explainable reference. However, there are several pitfalls in obtaining comparable and defensible reference values, see e.g., [8] and the application example in Sec. 5.3. In this respect, applying the PRB principle faces the same challenges as with GAMAB. In addition, PRB could be misused to offset risks from systems with little or no safety benefit.

In summary, all commonly known approaches feature limitations in the applicability to ADS. Hence, no single approach has yet been considered as gold standard.

5 Proposed approach

In the following, first, in Sec. 5.1, a new risk reference is introduced – the ENFLI (effectively no fleet incidents) approach – which can be used alternatively or in addition. Second, an overall framework is described in Sec. 5.2 which combines multiple approaches for the derivation of quantitative RAC. The framework is illustrated with an example in Sec. 5.3 and the findings are discussed in Sec. 5.4.

5.1 New risk reference: Effectively no fleet incidents (ENFLI)

The previously outlined approaches from the literature have in common that some *rate* of a critical event (e.g., accidents with fatalities per operating time) is used as a risk reference. Exemplary values, see e.g., [6], such as less than one fatal accident per ten million hours of ADS operation, are very small from the perspective of an individual vehicle user or affected non-user.

However, there is a serious implication if an ADS that just achieves such a target rate will be released in a vehicle fleet with a typical size for privately-owned vehicles. Simple calculations with a typical fleet size (e.g., $\geq 100\,000$ vehicles) and an expected average operating time of the ADS per vehicle over its lifetime (e.g., ≥ 1000 h) yield that the *expected value of the number* of critical events over the entire fleet is clearly larger than one. In practice, this can mean that it is expected that accidents will be caused by the fleet of ADS equipped vehicles over their lifetime. Given that serious consequences can result from even a single critical event, this raises doubts whether defensible RAC for an ADS could be derived with such rate-based criteria.

Therefore, the aim in the following is to make the probability of a critical event so small that in a realistic series application such an event is *effectively never expected to occur* over the lifetime of the product. The method is conceptually derived from industry practices for quantitative risk assessments that are used to evaluate potential field issues.

The criterion can be mathematically expressed with a probability or an expected value for the number of critical events over the product lifetime:

$$P(\text{Number of critical events} > 0) \ll 1 \text{ or } \mathbb{E}[\text{Number of critical events}] \ll 1. \quad (1)$$

Numerically there is no significant difference between the two formulae.³ We assume that

³Consider that on the one hand, the expected value for the number of critical events X reads $\mathbb{E}[X] = \sum_{k=1}^{\infty} k \cdot P(X = k)$. On the other hand, the probability of $X > 0$ events can be written as $P(X > 0) = \sum_{k=1}^{\infty} P(X = k)$. Thus, the only difference between the probability and expected value approach is that the probabilities $P(X = k)$ for $k > 1$ in the sum are weighted by k instead of 1. However, those probabilities should be very small given that the entire sum shall be $\ll 1$.

the expected value argumentation is easier to communicate, hence we use this in the sequel.

To calculate an expected value of a number of critical events, the event type needs to be defined. Preference is given to the number of **criticality normalised safety goal violations** (CNSGV) which is a leading measure, cf. Sec. 2. The term *criticality normalised* refers to a distinction according to the criticality of the safety goal as expressed by the ASIL rating in ISO 26262 [4]. In the proposal presented here, a difference in the ASIL rating is allowed to have an influence on the numerical values.

Much smaller than one is not a directly usable formula, so a more precise definition than $\ll 1$ must be introduced. Consistent with risk assessment experiences a range with a safety margin of about two orders of magnitude is proposed. Then, the main condition can be formulated as

$$\mathbb{E}[\text{CNSGV}] < E_{\text{limit}} \text{ where } E_{\text{limit}} \in [0.01, 0.05] . \quad (2)$$

While in each application an exact value will be assigned to E_{limit} , we consider it important to keep the range in mind. This also reflects the inherent uncertainty of the problem.

Typically, the quantification will be applied on a per safety goal basis. If multiple safety goals are relevant this might be critiqued as not conservative enough, but this is considered consistent with ISO 26262 [4]. As long as the number of safety goals is not too large, the margin implied in $E_{\text{limit}} \in [0.01, 0.05]$ is judged as being sufficient.

To compute $\mathbb{E}[\text{CNSGV}]$, an adequate fault model is required. Two standard cases are:

1. **Rate:** The RAC can be reasonably modelled as a per hour (or per another unit of time) failure rate r_{RAC} . This is useful when the failure phenomenon can continuously cause hazardous behaviour. Then $\mathbb{E}[\text{CNSGV}] = r_{\text{RAC}} \cdot RIF$ with some risk influencing factors combined in the variable RIF . This yields

$$r_{\text{RAC}} < E_{\text{limit}}/RIF . \quad (3)$$

2. **Probability on demand:** The RAC can be reasonably modelled as a per situation failure probability p_{RAC} . This is particularly useful when the failure phenomenon can only cause hazardous behaviour in a particular, “discrete” situation, e.g., a parking situation. Similarly, one obtains

$$p_{\text{RAC}} < E_{\text{limit}}/RIF . \quad (4)$$

Deriving risk influencing factors (combined in RIF) for the corresponding system and safety goal is an important part in the application of the ENFLI approach. Typical examples for risk influencing factors are listed below:

- **Number of vehicles:** Almost every risk scales with the number of vehicles in the fleet. Therefore, this parameter should contribute to the risk modelling which is a major feature of ENFLI compared to other approaches.
- **(Average) operating time of the function under consideration over the vehicle lifetime** either as a time value (failure modelling as a rate) or as a unitless expected number of situations (failure modelling as a probability on demand).

- **Normalisation factor derived from the ASIL classification parameters of risk:** E.g., equal to one order of magnitude per ASIL from 1.0 for ASIL D to 0.001 for ASIL A. Note that one should not inadvertently take the operating time of the function into account twice, explicitly and as part of an exposure parameter in the ASIL classification.

There might be qualitative arguments to relax the RAC number. For this, an expert judgement with a rationale is needed and the overall reduction should be restricted to two orders of magnitude. If such a rationale is available and with a choice of $E_{\text{limit}} = 0.01$ one obtains $\mathbb{E}[\text{CNSGV}] < 1$ from (2) which still implies that essentially no incident is expected, although no buffer is left. The following arguments may e.g., be considered to argue that no such buffer is needed:

- **Introduction of new technology:** The assumption here is that with the introduction of a new technology it is to some extent more easily accepted by society that “safety perfection” is less realistic than for mature technologies. This may cause some level of leniency towards the technology. A prerequisite is that if some incidents occur the causes are investigated, analysed, and removed.
- **Positive risk balance:** The assumption is that society is more willing to accept a risk if the introduction of the technology significantly reduces another risk. This is not the same as making the PRB a quantitative criterion of its own.

5.2 A generally applicable framework

The introduced ENFLI approach is particularly beneficial for innovative high-risk systems where little field experience of comparable systems is available and in case of an initial market introduction with a limited fleet size. However, there are cases in which other approaches can be used to derive RAC.

Thus, it is proposed to combine several approaches depending on certain criteria. The criteria and resulting approaches are visualized in a decision graph in Fig. 2. Three expert decisions are required as shown in the upper part:

1. Is one or more **quantitative RAC needed** at all or are qualitative criteria sufficient? The answer and rationale are part of the overall safety case.
2. Is an **unlimited market introduction possible** from a safety point of view, considering the available evidence?
3. Is there **sufficient data from comparable reference systems** available to derive risk acceptance criteria according to the GAMAB principle? Is a **demonstration of lower risk compared to human drivers** required by relevant regulations?

Depending on which path in the decision matrix is followed, some or all the activities shown in the lower part of Fig. 2 become relevant. Note that the risk reference derived from human accident statistics in the GAMAB (comparison to human driver) approach needs to be decomposed to reference values for the relevant hazardous behaviours.⁴

⁴See ISO 21448 [5, Annex C2.1] where an acceptance criterion is decomposed by dividing by the conditional probability of exposure to a scenario where the hazardous behaviour can lead to harm, $P(E|HB)$, the probability that the hazardous behaviour is not controllable, $P(C|E)$, and that the severity of the harm created matches the severity addressed with the acceptance criterion, $P(S|C)$.

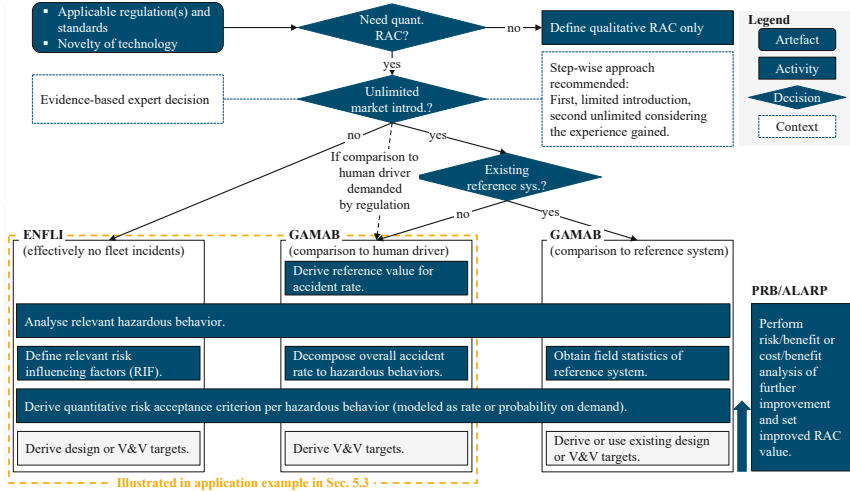


Figure 2: General framework to consider for the derivation of quantitative risk acceptance criteria for AD systems.

5.3 Application example

To illustrate the previously introduced risk references and associated activities, the ENFLI and the GAMAB (comparison to human driver) approach will be used to derive acceptance criteria for one hazardous behaviour of a hypothetical ADS.

The hypothetical system under consideration is an automated lane-keeping system (ALKS) for passenger cars which controls the longitudinal and lateral movement including lane change manoeuvres without further driver intervention, cf. [1]. The operational design domain (ODD) is limited to motorways with structurally separated driving directions and a travelling speed of up to 130 km/h.

Out of several safety goals relevant for the ALKS the example considers “collisions due to erroneous lateral guidance shall be avoided”. Parameters and assumed values relevant for the derivation of the RAC are given in Tab. 3.

For the ENFLI approach, the rate-based criterion (3) with $E_{\text{limit}} = 0.05$ yields

$$r_{\text{RAC,ENFLI}} < \frac{E_{\text{limit}}}{RIF} = \frac{0.05}{10\,000 \cdot 1000 \text{ h} \cdot 0.01} = 5 \times 10^{-7} / \text{h} . \quad (5)$$

The GAMAB (comparison to human driver) criterion is derived in three consecutive steps: First, a global reference value for the rate of fatal accidents by human drivers is required. This requires statistics on the number of accidents and the distance travelled for a comparable ODD, e.g., limited to passenger cars on motorways. We refer to the data sources and methodology applied in [8] for this example. Furthermore, a value for the distance between two accidents can be converted to a rate over time assuming an average travelling speed. This way, one obtains an average rate of $\approx 1.5 \times 10^{-7} / \text{h}$.

Although purely statistical fluctuation of this average can be additionally accounted for, the average is nonetheless dominated by non-rule conform drivers and includes older vehicles with lower technical standard. Unfortunately, there is no established approach or data basis on which a defensible margin for correcting these effects could be derived. Thus, to be on the safe side, a pessimistic safety margin of approximately two orders of magnitude is assumed which leads to a global reference value of $r_{\text{ref, global}} = 1 \times 10^{-9} / \text{h}$ for fatal accidents. This margin is to be regarded only as a proposal and other values might be defensible as well.

Second, the global reference value for fatal accidents from the first step is decomposed into individual rates for accidents related to specific hazardous behaviours. In this case, the share of accidents $\eta_{\text{specific, HB}}$ due to “erroneous lateral guidance” is estimated. This is done by estimating the share of relevant accidents where crossing a lane marking was the first event from the German in-depth accident study (GIDAS) [9].

Furthermore, the fault model which links the hazardous behaviour “erroneous lateral guidance” to the occurrence of a fatal accident is established. This is reflected in three conditional probabilities $P(E|HB)$, $P(C|E)$ and $P(S|C)$, see Tab. 3 for explanations and assumed values for the example.

Third, the decomposition of the global reference value $r_{\text{ref, global}}$ yields the following acceptance criterion for the hazardous behaviour “erroneous lateral guidance”:

$$r_{\text{RAC, GAMAB}} < \frac{r_{\text{ref, global}} \cdot \eta_{\text{specific, HB}}}{P(E|HB) \cdot P(C|E) \cdot P(S|C)} = \frac{1 \times 10^{-9} / \text{h} \cdot 0.5}{0.01} = 5 \times 10^{-7} / \text{h} . \quad (6)$$

Therefore, for a limited market introduction with only 10 000 vehicles and 1000 h of operating time, the RAC for the considered hazardous behaviour derived with the ENFLI approach is the same as when applying the GAMAB principle.

However, while the GAMAB approach is independent of the fleet size, the results from the ENFLI approach scale linearly with the fleet size. In practice, the latter cannot grow unboundedly. This motivates to estimate the order of magnitude of a limit value $\bar{r}_{\text{RAC, ENFLI}}$. To this end, we consider a hypothetical ADS with the following properties:

- The safety goal can be violated continuously over time, thus the RAC can be reasonably modelled as a per hour rate.
- A deployment to a large-scale platform is intended and thus a fleet size of $\leq 1 \times 10^7$ vehicles is assumed. Fleets beyond that size are very rare and the size increase is not by an order of magnitude. These rare cases are hence covered by the margin implied in $E_{\text{limit}} \in [0.01, 0.05]$.
- Concerning the contributions to the risk influencing factor, an average vehicle lifetime of 10 000 h, 100 % operating time of the function and a normalization factor of 1.0 based on the maximum ASIL classification are assumed.

Then, inserting into (5) yields a limit value of

$$\bar{r}_{\text{RAC, ENFLI}} < \frac{E_{\text{limit}}}{RIF} = \frac{0.01}{1 \times 10^7 \cdot 10\,000 \text{ h} \cdot 1.0} = 1 \times 10^{-13} / \text{h} . \quad (7)$$

This value might look very conservative but comparing it with the level achieved by mature safety related automotive systems, at least with respect to fatalities, this is not far from reality.

Table 3: Relevant parameters and assumed values for application example in Sec. 5.3.

Parameter	Value	Rationale
Contributions to risk influencing factor (RIF) in ENFLI approach		
Number of vehicles	10 000	Initial market introduction of a limited amount
(Average) operating time of the ALKS over vehicle lifetime	1000 h	Average vehicle lifetime of 10 000 h and thereof 10 % ALKS operation given ODD limitation
Normalization factor derived from the ASIL classification	0.01	Based on classification of “erroneous lateral guidance” as ASIL B due to low severity in motorway situations as determined from accident statistics.
Risk reference derived from accident statistics		
Global reference value for the rate of fatal accidents on motorways involving passenger cars	$r_{\text{ref, global}} = 1 \times 10^{-9} / \text{h}$	• Number of fatal accidents with passenger cars on German motorways in 2015-2019 from DESTATIS [10] • Annual distances travelled by passenger cars on motorways as in [8] • Conversion to rate per hour assuming 110 km/h average travelling speed • Additional safety margin of two orders of magnitude
Share of accidents due to specific hazardous behaviour	$\eta_{\text{specific, HB}} = 0.5$	Based on analysis of accidents described in GIDAS [9] triggered by passenger cars with electronic stability control on motorways.
Conditional probability of exposure to a scenario where the hazardous behaviour can lead to harm	$P(E HB) = 1.0$	Practically in all scenarios for “erroneous lateral guidance”
Probability that the hazardous behaviour is not controllable	$P(C E) = 1.0$	Pessimistic choice
Probability that the injury severity of harm due to an uncontrolled hazardous behaviour is fatal	$P(S3 C) = 0.01$	Based on analysis of accidents described in GIDAS [9] where crossing a lane marking was the first event.

5.4 Discussion

While this outcome is not surprising given the introductory remarks in Sec. 5.1, it raises again the question for which purposes quantitative RAC can be useful. Concerning the purpose of designing a system (architecture), it makes sense to consider targets derived according to the ENFLI approach to develop a sufficiently strong safety architecture that allows for a scaled-up market introduction. However, concerning the purpose of deriving V&V targets, it will become extremely challenging to demonstrate that an RAC nearing $\bar{r}_{\text{RAC,ENFLI}}$ is achieved before market introduction.

Thus, it is important to differentiate between a quantified risk value as an acceptance criterion and quantified confidence obtained from V&V results that the former is achieved. Even if there is a gap between the statistical confidence obtained from V&V results before market introduction and the RAC, qualitative measures can prevent this gap from materialising as excess risk.

One qualitative argument with a quantifiable risk-limiting effect is scrutiny in the observation of field incidents. If the field observation and control of operation are effective, the excess risk is limited to singular critical events. Note that the choice of (normalised) safety goal violations as the type of critical event used in the RAC definition has an additional risk-limiting influence.

Clearly, the gap between RAC values derived using the ENFLI approach and quantified confidence obtained from achievable V&V targets can be made small by initially introducing only a limited quantity of the system in the market. The experience gained with this fleet can inform a subsequent larger introduction.

6 Summary

Quantitative RAC are a subset of a more holistic set of acceptance criteria and while they are clearly useful for some purposes, e.g., evaluating a system architecture or deriving V&V targets, other purposes are better addressable by qualitative criteria and measures. Moreover, quantitative criteria are well suited to address some risks, e.g. stemming from functional insufficiencies, but are not applicable to others, e.g., security attacks.

Known approaches for defining quantitative RAC are either too generic to yield defensible criteria for an ADS (e.g., MEM), or their application raises additional questions for which however sufficiently detailed data and models is lacking (e.g., GAMAB, PRB).

We identify that there is a conceptual issue that such approaches consider the rate of critical events normalised per vehicle whereas, in practice, a manufacturer releases an entire fleet of vehicles. Motivated by the fact that serious consequences can result from even a single critical event in the field, this paper proposes a new risk reference targeting effectively no fleet incidents (ENFLI).

We present an illustrative example, which shows that the RAC values obtained with the ENFLI approach quickly outgrow those derived from a comparison against the accident rate of human drivers (GAMAB approach). On the one hand, this is a desirable property of an RAC as it helps designing a sufficiently strong system safety architecture that allows for a scaled-up market introduction. On the other hand, a gap can occur before release between the RAC value and quantifiable confidence obtained from V&V results. However, as we discuss, the size of such a gap does not imply a proportionally

5.4 Discussion

While this outcome is not surprising given the introductory remarks in Sec. 5.1, it raises again the question for which purposes quantitative RAC can be useful. Concerning the purpose of designing a system (architecture), it makes sense to consider targets derived according to the ENFLI approach to develop a sufficiently strong safety architecture that allows for a scaled-up market introduction. However, concerning the purpose of deriving V&V targets, it will become extremely challenging to demonstrate that an RAC nearing $\bar{r}_{\text{RAC,ENFLI}}$ is achieved before market introduction.

Thus, it is important to differentiate between a quantified risk value as an acceptance criterion and quantified confidence obtained from V&V results that the former is achieved. Even if there is a gap between the statistical confidence obtained from V&V results before market introduction and the RAC, qualitative measures can prevent this gap from materialising as excess risk.

One qualitative argument with a quantifiable risk-limiting effect is scrutiny in the observation of field incidents. If the field observation and control of operation are effective, the excess risk is limited to singular critical events. Note that the choice of (normalised) safety goal violations as the type of critical event used in the RAC definition has an additional risk-limiting influence.

Clearly, the gap between RAC values derived using the ENFLI approach and quantified confidence obtained from achievable V&V targets can be made small by initially introducing only a limited quantity of the system in the market. The experience gained with this fleet can inform a subsequent larger introduction.

6 Summary

Quantitative RAC are a subset of a more holistic set of acceptance criteria and while they are clearly useful for some purposes, e.g., evaluating a system architecture or deriving V&V targets, other purposes are better addressable by qualitative criteria and measures. Moreover, quantitative criteria are well suited to address some risks, e.g. stemming from functional insufficiencies, but are not applicable to others, e.g., security attacks.

Known approaches for defining quantitative RAC are either too generic to yield defensible criteria for an ADS (e.g., MEM), or their application raises additional questions for which however sufficiently detailed data and models is lacking (e.g., GAMAB , PRB).

We identify that there is a conceptual issue that such approaches consider the rate of critical events normalised per vehicle whereas, in practice, a manufacturer releases an entire fleet of vehicles. Motivated by the fact that serious consequences can result from even a single critical event in the field, this paper proposes a new risk reference targeting effectively no fleet incidents (ENFLI).

We present an illustrative example, which shows that the RAC values obtained with the ENFLI approach quickly outgrow those derived from a comparison against the accident rate of human drivers (GAMAB approach). On the one hand, this is a desirable property of an RAC as it helps designing a sufficiently strong system safety architecture that allows for a scaled-up market introduction. On the other hand, a gap can occur before release between the RAC value and quantifiable confidence obtained from V&V results. However, as we discuss, the size of such a gap does not imply a proportionally

growing excess risk in the field when field observation is effective.

References

- [1] United Nations Economic Commission for Europe, “01series of amendments to UN Regulation No. 157: Uniform provisions concerning the approval of vehicles with regard to Automated Lane Keeping,” 2022.
- [2] European Commission, “Commission Implementing Regulation (EU) 2022/1426: Uniform procedures and technical specifications for the type-approval of the automated driving system (ADS) of fully automated vehicles,” in *Official Journal of the European Union (Vol. 65, L 221)*, pp. 1–64, 2022.
- [3] Bundesministerium für Digitales und Verkehr, “Verordnung zur Regelung des Betriebs von Kraftfahrzeugen mit automatisierter und autonomer Fahrfunktion und zur Änderung straßenverkehrsrechtlicher Vorschriften,” in *Bundesgesetzblatt (Jg. 2022, Teil I, Nr. 22)*, pp. 986–1010, 2022.
- [4] International Organization for Standardization, “ISO 26262:2018 Road vehicles – functional safety” Geneva, Switzerland, 2018-12.
- [5] International Organization for Standardization, “ISO 21448:2022 Road vehicles – safety of the intended functionality,” Geneva, Switzerland, 2022-06.
- [6] C. Amersbach, T. Ruppert, N. Hebgen and H. Winner, “Macroscopic Safety Requirements for Highly Automated Driving in Urban Environments,” in *13. Graz Symposium Virtual Vehicle (GSVF)*, Graz, 2020.
- [7] L. Putze, L. Westhofen, T. Koopmann, E. Böde, and C. Neurohr, “On Quantification for SOTIF Validation of Automated Driving Systems,” *2023 IEEE Intelligent Vehicles Symposium (IV)*, Anchorage, 2023.
- [8] F. Fahrenkrog, L. Drees, and F. Raisch, “Implications of the positive risk balance on the development of automated driving,” in *Traffic Injury Prevention (24sup1)*, pp. 124–130, 2023.
- [9] D. Otte, J. Michael, and H. Carl, “Scientific approach and methodology of a new in-depth investigation study in Germany called GIDAS,” in *Proceedings of the International Technical Conference on the Enhanced Safety of Vehicles*, 2003. Parameters used in the example were derived from the data available in GIDAS until 12/2019.
- [10] DESTATIS (Federal Statistical Office of Germany), “Verkehrsunfälle in Deutschland,” https://www.destatis.de/DE/Themen/Gesellschaft-Umwelt/Verkehrsunfaelle/_inhalt.html#sprg478574.

SURE-Val: Safe Urban Relevance Extension and Validation

Kai Storms^{*†}, Ken Mori^{*†} und Steven Peters[†]

Abstract: To evaluate perception components of an automated driving system, it is necessary to define the relevant objects. While the urban domain is popular among perception datasets, relevance is insufficiently specified for this domain. Therefore, this work adopts an existing method to define relevance in the highway domain and expands it to the urban domain. While different conceptualizations and definitions of relevance are present in literature, there is a lack of methods to validate these definitions. Therefore, this work presents a novel relevance validation method leveraging a motion prediction component. The validation leverages the idea that removing irrelevant objects should not influence a prediction component which reflects human driving behavior. The influence on the prediction is quantified by considering the statistical distribution of prediction performance across a large-scale dataset. The validation procedure is verified using criteria specifically designed to exclude relevant objects. The validation method is successfully applied to the relevance criteria from this work, thus supporting their validity.

Keywords: Automated Driving, Perception, Relevance, Safety

1 Introduction

Automated driving (AD) is currently viewed as a key emerging technology. Among the benefits which are attributed to AD are increased comfort, availability of mobility and foremost an increase in safety for the traffic environment [11]. However, safety assurance of AD remains a challenge. Whilst first SAE level 3 AD system (AD) are already available [2] for parts of the highway domain, a current focus in research is the urban domain [13].

One method supporting safety assurance by offering potential benefits regarding the testing effort is modular decomposition [4]. Under this method, all modules of a classic Sense-Plan-Act architecture and its derivatives, such as the perception module, need to be evaluated individually [38]. For this individual evaluation, it is imperative to have knowledge about what is relevant to the subject module, both for completeness and efficiency.

Therefore, in this paper we will consider relevance for the perception module and how it can be evaluated. Current approaches face a variety of problems. Most methods lack generality because they are specific to either one module or a limited set of scenarios. These methods cannot be applied to other modules or scenarios without requiring major redesign efforts.

^{*}contributed equally

[†]Institute of Automotive Engineering (FZD) at Technical University of Darmstadt, 64287 Darmstadt (e-mail: firstname.lastname@tu-darmstadt.de)

Furthermore, a lack of consistency between relevance results must be noted among different methods. One reason is the lack of validation in previous methods. This situation is exacerbated by the fact that the current state of the art does not provide any methodology to validate relevance criteria.

This paper will expand on the previous work of Mori & Storms [30], henceforth denoted as Structured Analysis for Conservative Relevance Estimation in Driving context (SACRED). Firstly, the Safe Urban Relevance Extension (SURE) is introduced. As main contribution, we present a novel methodology that enables the evaluation of validity for a given relevance selection method. This methodology is applied and verified using the expanded method for the urban domain.

2 Related Work

This section covers related works with respect to concepts of relevance as well as their respective validation.

2.1 Relevance

Previous conceptualizations of relevance can be broadly categorized in heuristic approaches, formal approaches and concepts based on a downstream task.

Heuristic approaches typically implicitly define relevance by excluding certain objects based on simple criteria. One application is the datasets used to develop and test perception functions. Here, thresholds on distance [42] or number of points from lidar and radar [9] are applied during annotation. Heuristics are also applicable when defining perception metrics [21]. Criteria such as distance [9], height in camera image plane and occlusion [18] are used by dataset metrics. Other metric proposals in literature follow similar approaches such as leveraging the distance to the ego vehicle [27] or the ego trajectory [8]. Similar ideas are found in neural path planners where relevance is implicitly defined by the network input. Commonly, inputs are restricted to a certain geometric region [6, 19, 34, 40] or a limited number of objects [1, 12, 17, 23, 44].

Formal approaches provide an explicit consideration of relevance based on requirements. Typically, relevance is related to safe behavior of the vehicles by considering reachability or formal planners [21]. Reachability leverages kinematic constraints to define potential collision objects as relevant [3, 43]. Other work leverages formal planners either from preexisting work [45] or by directly specifying context-dependent behavioral requirements [30, 37, 41].

In order to avoid manual specification of relevance, the planning Kullback-Leibler divergence (PKL) [35] and following work [20, 37] propose to leverage neural planners. Here, relevance is conceptualized and quantified as magnitude of the effect of an object on a downstream planner implementation [35]. However, the validity of this approach is limited to the specific implementation of the planning algorithm [36].

Overall, various approaches for the definition of relevance are proposed and reach different conclusions. Notably, there is currently no reconciliation of formal and downstream implementation based approaches.

2.2 Relevance Validation

While different approaches have been suggested to conceptualize object relevance, the validation of these concepts is currently underexplored. Perception datasets such as [18, 42, 47] do not consider the validity of the perception metrics with respect to safety. Other datasets discuss metrics with respect to the weighting of different attributes [28] or validate the ability of the metric to produce a ranking between different types of detectors [9].

Some analytic approaches rely exclusively on plausibilization by visualizing selected scenarios without further validation [37, 43, 45]. In other cases, no further attempts to verify or validate the results are made [14].

The usage of planners in a closed loop simulation has been proposed [43] and also executed for safety aware prediction metrics [24]. However, incorporating a planner into the pipeline incurs the potential of errors in the planner which questions the validity [29]. The results of the PKL metric are verified and validated in two different ways. A first verification step argues plausibility by showing that the metric is sensitive to distance and velocity as intuitively salient features. Furthermore, a human evaluation is conducted by Amazon Mechanical Turk workers, showing an 80% preference for the PKL metric over the nuScenes detection score (NDS).

Overall, current validation approaches for validating relevance criteria are lacking. Furthermore, the suggested approaches generally focus on relative ranking and therefore do not provide acceptance criteria.

3 Methodology

We suggest the following approach for defining and validating relevance in this work. First, a formal analytic method of relevance is selected to facilitate interpretability. Next, a neural trajectory prediction model is utilized to validate the results of the formal model. It will later be shown that neural networks are suitable for validation if large-scale datasets are leveraged to consider uncertainties in the prediction outcome.

As formal relevance method, the prior SACRED method [30] is selected. This method provides a conservative estimation aiming to provide a complete set of relevant objects. The property of completeness will also be focus of the later validation procedure. In addition, the relevance method offers the benefit of few requirements regarding environment knowledge.

4 Application of Relevance Selection Method

This section follows the SACRED relevance method defined in [30]. First, a minimum specification is provided. Next, interactions are decomposed into multiple scenarios for each of which relevance criteria are formulated.

4.1 Specification

In order to apply the relevance methodology, it is first necessary to specify the system and the use case.

The system specification follows SACRED [30]. A Sense-Plan-Act architecture is assumed with an object list as interface between perception and planning. The output trajectories of the planner must generally adhere to physical and legal requirements. This work focuses on the task of collision safety. In order to fulfill these requirements, the system provides guarantees regarding its capabilities. After a specified reaction time, the system is required to act in accordance with the requirements. In order to do so, the system is capable of providing a minimum guaranteed braking deceleration as well as a minimum guaranteed acceleration.

Currently, testing of perception is mainly performed on datasets such as [9] which mainly consider urban environments. Therefore, the use case is selected to be the application in the urban domain. This use case serves to expand the ideas from the highway domain to more complex interactions.

4.2 Decomposition

The next step in determining relevance is the decomposition of the use case into functional scenarios. As in SACRED [30], the distinction is made using polar coordinates for the relative velocities, distinguishing radial and tangential scenarios (first letter of abbreviation). Depending upon whether the ego is moving towards or away from the object of interest (OOI)(second letter of abbreviation) and whether the OOI is moving towards or away from the ego (third letter of abbreviation), the four radial scenarios R.TA, R.AT, R.TT and R.AA are distinguished. The tangential scenarios are distinguished depending on whether the OOI is moving towards or away from the ego vehicle. This is described by the two scenarios T.XT and T.XA.

The results for the R.TA, R.AT, R.AA and T.XA scenario are identical to SACRED [30] and are thus not repeated here. However, the R.TT and the T.XT scenario require additional considerations to consider the urban environment. In order to operate in an urban environment a vehicle needs to be able to pass static obstacles by entering the opposite lane and to traverse intersections. Overtaking dynamic objects is excluded to limit complexity, since it is not required to fulfill a driving task. These are considered by modifying the R.TT and the T.XT scenario respectively. Each of the scenarios requiring modification to the relevance estimation is discussed in the following subsections.

4.3 R.TT'

First, the R.TT scenario is applied as in SACRED [30] without modification. For all objects not considered relevant by the R.TT scenario, a modified R.TT scenario is used. This scenario, shown in Fig. 1, hereafter called R.TT', takes into account the case of passing a static object.

Variables follow the notation from SACRED [30] with three optional indices. The first index identifies the object, with 1, 2 and 3 denoting the ego vehicle, opposing vehicle and static object respectively. The second index gives the frame of reference, either r for radial or t for tangential. The third index denotes the state.

We extend the SACRED method [30] which only considers pairwise interactions. First, the static object is considered as OOI. Since the OOI is static, the R.TT and R.TA are equivalent. In order to distinguish the R.TT' scenario, the static object has to be relevant

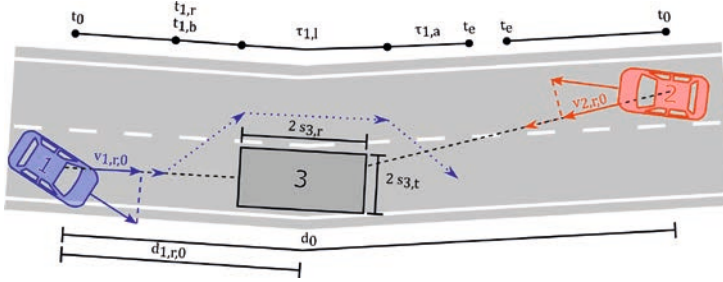


Figure 1: Sequence and variables for the R.TT' scenario.

according to either scenario. Arbitrarily selecting R.TA requires the distance to the static object $d_{1,r,0}$ to be less than the sum of reaction and breaking distance:

$$0 < d_{1,r,\min} = d_{1,r,0} - s_1 - s_{3,r} - v_{1,r,0} t_{1,r} + \frac{1}{2} a_{\max} t_{1,r}^2 - \frac{(v_{1,r,0} + t_{1,r} a_{\max})^2}{2a_{1,r,b}} \quad (1)$$

The R.TT' scenario is applied if all following conditions for the initial positions $r_{i,0}$ and velocities $v_{i,0}$ are fulfilled:

- Ego moving towards static object: $(\vec{r}_{3,0} - \vec{r}_{1,0}) \cdot \vec{v}_{1,0} > 0$
- OOI moving towards static object: $(\vec{r}_{3,0} - \vec{r}_{2,0}) \cdot \vec{v}_{2,0} > 0$
- OOI located behind static object: $(\vec{r}_{3,0} - \vec{r}_{1,0}) \cdot (\vec{r}_{3,0} - \vec{r}_{2,0}) < 0$

The scenario is described by the following considerations, according to worst case assumptions. Within the reaction time, the ego vehicle brakes. This results in a distance covered and velocity as defined in (2) and (3). The end time of the braking maneuver $t_{1,b}$ in (4) is equal to the reaction time or the time to standstill if less.

$$d_{1,b} = v_{1,r,0} t_{1,b} - \frac{1}{2} a_{\max} t_{1,b}^2 \quad (2)$$

$$v_{1,r,b} = v_{1,r,0} - a_{\max} t_{1,b} \quad (3)$$

$$t_{1,b} = \min \left\{ t_{1,r}, \frac{v_{1,r,0}}{a_{\max}} \right\} \quad (4)$$

After the reaction time, the passing maneuver is initiated with a lateral movement to the opposite lane, followed by accelerating with the guaranteed acceleration. After passing the static object, the ego vehicle performs a lateral movement back to its initial lane, thus concluding its maneuver. The durations of the lateral lane change $\tau_{1,l}$ and acceleration $\tau_{1,a}$ are defined in (5) and (7). The final velocity after accelerating $v_{1,a}$ is given by (6).

$$\tau_{1,l} = 2 \sqrt{\frac{s_{3,t} + s_1}{a_{1,g}}} \quad (5)$$

$$v_{1,r,a} = v_{1,r,b} + a_{1,g}\tau_{1,a} \quad (6)$$

$$\tau_{1,a} = \frac{-v_{1,r,b} + \sqrt{2a_{1,g}(2s_1 + 2s_{3,r}) + v_{1,r,b}^2}}{a_{1,g}} \quad (7)$$

The final distance the ego vehicle traverses during the scenario is defined in (8) as sum of the individual maneuvers.

$$d_{1,e} = d_{1,b} + v_{1,r,b}\tau_{1,l} + d_{1,a} + v_{1,r,a}\tau_{1,l} \quad (8)$$

Throughout the scenario, the opposing vehicle performs the maximum acceleration towards the ego, yielding (9) and (10) with (11).

$$d_{2,e} = v_{2,r,0}t_e + \frac{1}{2}a_{\max}t_e^2 \quad (9)$$

$$v_{2,r,e} = v_{2,r,0} + a_{\max}t_e \quad (10)$$

$$t_e = t_{1,r} + 2\tau_{1,l} + \tau_{1,a} \quad (11)$$

Given the resulting distances and velocities of the defined events, the relevance of the opposing vehicle can then be determined by using them as the initial values for R.TT from SACRED [30] as defined in equation (12).

$$0 < d_{\min} = (d_0 - d_{1,e} - d_{2,e}) - s_1 - s_2 - v_{1,r,a} t_{1,r} - \frac{1}{2}a_{\max}t_{1,r}^2 - \frac{(v_{1,r,a} + t_{1,r}a_{\max})^2}{2a_{1,r,b}} - v_{2,r,e}t_{1,b} - \frac{1}{2}a_{\max}t_{1,b} \quad (12)$$

Exemplary scenarios yield distances in the hundreds of meters, well beyond typical dataset annotation ranges. In addition, passing only needs to be considered given an ego intention. Since this information is not available on datasets, this scenario is neglected for the practical implementation.

4.4 T.XT'

This scenario covers tangential movements such as merging and intersections. The merging covered in SACRED [30] is extended to intersections encountered in the urban domain.

Within an intersection, there are two distinct maneuvers available to the ego vehicle as shown in Fig. 2. It can either merely pass through (A) or turn onto the path of another vehicle travelling (B). The requirement in either case is not to impede another vehicle. The latter case represents the worst case since ego vehicle must not only leave the intersection, but additionally accelerate to an adequate speed.

Thus, the worst case in an intersection is adequately described by the T.XT scenario from SACRED [30]. The main difference is that the lateral distance to the travelling direction of the other vehicle can be arbitrarily large unlike for the case of the highway domain. Therefore, the ego vehicle has the opportunity to avoid entering the intersection by braking in time. The corresponding braking distance is:

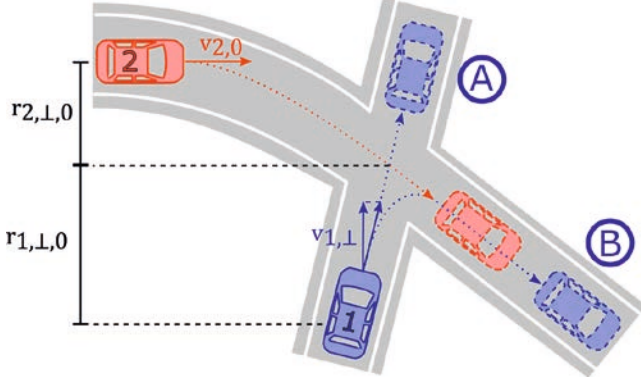


Figure 2: Worst-case intersection and variables in the T.XT' scenario.

$$s_{1,\perp,b} = v_{1,\perp,0}t_{1,r} + \frac{v_{1,\perp,0}^2}{2a_{1,\perp,b}} \quad (13)$$

Additionally, we assume the location of the intersection to be unknown. Therefore, the intersection may be closer to the ego than the heading of the other vehicle. In the worst case, the other vehicle would perform a maximum lateral acceleration towards the ego vehicle to reach the intersection. The corresponding lateral distance traveled with the corresponding braking time of the ego vehicle is:

$$d_{2,\perp} = \frac{1}{2}a_{\max}t_{1,b}^2 \quad (14)$$

$$t_{1,b} = t_{1,r} + \frac{v_{1,\perp,0} + t_{1,r}a_{\max}}{a_{1,b}} \quad (15)$$

Considering these two components, the maximum lateral distance for which the merging scenario from SACRED [30] is considered is limited to:

$$r_{1,\perp,0} < v_{1,\perp,0}t_r + \sqrt{\frac{v_{1,\perp,0}^2}{2a_{1,\perp,b}}} + \frac{1}{2}a_{\max} \left[t_{1,r} + \frac{v_{1,\perp,0} + t_{1,r}a_{\max}}{a_{1,b}} \right]^2 \quad (16)$$

If this equation is fulfilled, the merging scenario from SACRED [30] is applied. Otherwise, the respective object is considered irrelevant according to this scenario.

5 Validation

This section discusses the validation of the results presented thus far. First, some preliminary discussion relating to prior work is shown. Next, the approach of this work is presented. Finally, the approach is applied and verified.

5.1 Preliminaries

As noted in previous sections, the validation of relevance criteria has received limited attention with no generally accepted methodology available. Previous relevance concepts either take an argumentative approach or include a concrete downstream planning task. In addition, the notion of a human baseline is present in one work [35]. We believe that reconciliation of these complimentary approaches is required in order to argue validity.

One basis for the proposed validation approach is the human baseline. This idea has previously been applied for accident rates of human driving performance [25,26,33] as well as for perception performance [39]. Works incorporating downstream planners implicitly also include this concept. However, planning goals may include additional aspects other than imitating human behavior [6] such as infractions, mission goals [16] and comfort [10]. Therefore, these goals are ambiguous and lack consistent evaluation metrics [10].

5.2 Proposed Method

We conceptualize relevance validation by considering the human behavior as baseline. Ideally, the influence of removing objects considered irrelevant might be studied in driving simulators to directly quantify the behavioral changes in human subjects. To avoid the substantial costs such an approach incurs, we propose to approximate human driving behavior with a motion prediction algorithm. This approach follows the approach proposed in [30] which is adapted from [35]. A motion prediction algorithm has several advantages over using a path planning algorithm. The objective of the motion prediction is unambiguous and can be evaluated in an open-loop setting with real-world data without relying on a simulation environment. This work uses the predictor as proxy for human behavior rather than as component to create a full AD pipeline as in [35]. Additionally, this work explicitly quantifies the prediction performance as opposed to prior work which explicitly or implicitly assumes the planner to be valid [24,35].

The validation procedure consists of running the prediction network with different inputs. Predictions are calculated multiple times for each input to account for potential uncertainties and non-deterministic elements of the prediction. The validation procedure only considers the empirical cumulative distribution function (ECDF) of prediction errors across a large-scale dataset to avoid local performance issues and effects of non-deterministic predictions. A relevance criterion is considered valid if the prediction error distribution remains unchanged between the original input and the filtered input. Whether two distributions are identical is evaluated using the Cramer-von Mises test for two empirical distributions [5]. This test provides a confidence value, based on distribution similarity and sample size, for which an acceptance criterion is required.

To study the utility of the proposed validation method, an implementation on the nuScenes dataset [9] is performed. All experimental results are reported for the standard val split. The prediction network is selected from the nuScenes motion prediction leaderboard [32] among entries with an open implementation. We select the PGP algorithm [31] as implemented by [15] using the default settings. The implementation filters inputs by location in a square region from [-20 m, 80 m] in longitudinal and [-50 m, 50 m] in lateral direction. This filtering is maintained for all experimental conditions with the relevance filtering criteria only being additionally applied. The prediction errors are evaluated using the average distance error (ADE) metric for the top 10 trajectories.

5.3 Results

The prediction network is run ten times each with four different inputs. Firstly, multiple prediction runs with all inputs abbreviated as ‘A’ are compared to each other to determine the prediction noise. The ECDF of this noise is depicted in Fig. 3, showing that the detector exhibits significant noise averaging at 0.33 m. The error distribution is depicted as ECDF averaging at 0.96 m. Noise also affects the error distributions. Therefore, different runs using identical inputs with all objects (A-A) are compared using statistical tests for equality of distribution. The results are displayed in Fig. 4 as box plot of p-values for different runs. Since the distributions are similar despite the presence of noise, the p-values are high. This provides a reference for the p-values when accounting for noise.

To verify the validation procedure, two artificial verification inputs abbreviated with ‘RV’ and ‘RV2’ are constructed and compared with the original input. The first is simply deleting all objects in a scene. The second is to delete all vehicles which are within 2 m from the heading axis of the ego vehicle. Both inputs are constructed to be implausible relevance criteria. The latter filters out 5% of the objects within the heuristic input region of the prediction network, thus filtering out fewer objects than the relevance criteria developed in this work. The error distributions in Fig. 3 exhibit general visual similarity, all being larger than the prediction noise. Nevertheless, the enlarged image indicates that the verification inputs differ from the case using all inputs. Testing for equality of distributions between verification and all inputs (A-RV and A-RV2) in Fig. 4 shows p-values which are orders of magnitude smaller than for A-A. The large spread in p-values is a consequence of prediction noise. Average p-values are below the exemplary threshold value of 0.005 [7]. This indicates a high confidence that the error distributions are different for the verification input and all inputs. Since the prediction is impacted by the verification input, the verification criteria are successfully identified as invalid.

Similar to the verification inputs, the prediction is performed on the input filtered according to the relevance criteria of this work abbreviated as ‘R’. The relevance criteria filter out 10% of the objects included within the heuristic input region of the prediction network. Fig. 3 shows that even with an enlarged image, the distributions appear visually similar. The results of testing for equality of all and of relevant inputs (A-R) are shown in Fig. 4. The p-values are similar to A-A, indicating that any dissimilarity in error distributions is within the range caused by noise. Since the prediction performance is unaffected, the relevance criteria of this work are not falsified. Accordingly, the validation results support the relevance criteria from this work.

6 Discussion

This section discusses the results of both the validation method proposed as well as the validation of the extended urban relevance model.

6.1 Validation Method

The proposed validation method relies on a motion prediction network. Compared to PKL [35], fewer requirements are imposed upon the motion prediction. Discrete output trajectories and non-deterministic implementations are applicable without modification.

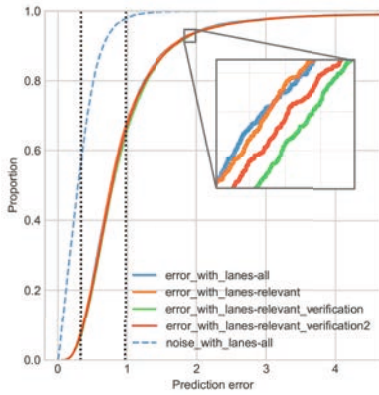


Figure 3: ECDF prediction error as ADE for top 10 trajectories for different inputs and noise for multiple runs with same inputs.

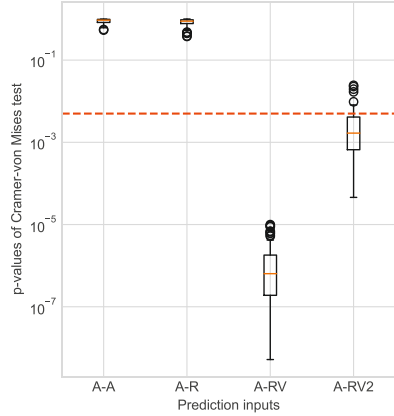


Figure 4: Boxplot of p-values testing for equality of distribution for pairwise combinations. Exemplary threshold indicated by red dashed line.

The method explicitly considers the global performance of the prediction, showing the error to be in the order of 1 m. Notably, the global performance is robust to local noise which may appear in the prediction for a single scenario even with identical inputs. We suspect the inherent uncertainty of future trajectories causes stochastic behavior in the predictions. Since the underlying trajectory distribution is unknown, the prediction performance can only be assessed over a large number of scenarios. Conversely, this means that the local prediction performance in a single scenario is not meaningful since it cannot be disentangled from the stochastic component. Therefore, the local stability of the prediction as used by PKL [35] is not suited as relevance measure for some prediction networks. This is visible in the noise results as exhibited by a non-deterministic prediction implementation as used in this work.

The validation method of this work is successful in falsifying the verification input. The global performance of the prediction deteriorates, which agrees with the intuition that the vehicle directly in front is relevant to the vehicle. This verification succeeds despite the fact that only 5% of the objects are filtered out, as opposed to the 10% for the relevance criteria. However, further study is required to understand the impact of thresholds for p-values. In addition, further research is required to determine to what degree invalid samples can be resolved in large datasets. However, higher sensitivity of the validation may be possible if subsets of specific scenarios are considered individually.

6.2 Extended Urban Relevance Model

Generally, we first develop analytic relevance criteria which are then reconciled with a prediction component at validation. We consider this approach to have the following ad-

vantages. Firstly, no complex neural networks are required when applying the relevance criteria. This is beneficial with regards to implementation effort [48] and computational resources. Additionally, neural networks are prone to failures due to their lack of robustness [22, 46] which may impact directly determining relevance. However, the proposed approach is robust to these failures since the distribution across many samples is considered. Additionally, the analytic relevance criteria are fully interpretable.

These criteria are developed by applying an existing method to derive relevance. The method is successfully applied to extend the results to the urban domain. Applying the proposed validation method on a large-scale dataset shows that the prediction performance is unaffected by the relevance criteria proposed in this work. The validation is currently restricted to a single prediction network and one dataset. For this case, the validity of the relevance criteria is supported. However, only 10% of objects can be excluded from the data based on relevance. It is currently not known if the specificity of the relevance model can be improved while maintaining validity.

7 Conclusion and Outlook

This work presents SURE-Val, the extension of a recent method to derive relevance to the urban domain. For this purpose, the methodology is applied to the intersection and the passing scenario. Additionally, a novel method for validating analytic relevance criteria using a motion prediction is introduced. The relevance criteria and the validation method are applied on a public dataset using an exemplary motion prediction component. The verification of the validation procedure shows that the validation method itself is feasible. At the same time, the validation results support the relevance criteria obtained in this work. Additionally, the analytic relevance criteria are computationally efficient, interpretable and agnostic with regards to the implementation of downstream components. We hope that both the relevance criteria as well as the validation method can serve as reference to consider these aspects more explicitly in the future.

For future research, it is desirable to further explore the limitations of the relevance criteria and their validation. One aspect is the impact of the prediction architecture on relevance. Another aspect is the consideration of datasets including more varied scenarios. Further validation on such datasets which include unusual and dangerous driving situations is required to gain confidence in relevance criteria in general.

8 Acknowledgement

The research leading to these results is funded by the German Federal Ministry for Economic Affairs and Climate Action within the project VVM - Verification Validation Methods under grant number 19A19002S, as well as by the German Federal Ministry of Education and Research within the project VIVID with the grant number 16ME0173 based on a decision of the Deutscher Bundestag. The authors would like to thank the consortia for the successful cooperation.

References

- [1] Zero-shot deep reinforcement learning driving policy transfer for autonomous vehicles based on robust control. *ArXiv*, 1812.03, 2018.
- [2] Mercedes-Benz AG. Mercedes-benz drive pilot. <https://www.mercedes-benz.de/passengercars/technology/drive-pilot.html>, last accessed 17.07.2023.
- [3] Matthias Althoff. *Reachability Analysis and its Application to the Safety Assessment of Autonomous Cars*. Dissertation, Technische Universität München, München, 2010.
- [4] Christian Amersbach. *Functional Decomposition Approach: Reducing the Safety Validation Effort for Highly Automated Driving*. Dissertation, Technische Universität Darmstadt, Darmstadt, 26.09.2019.
- [5] T. W. Anderson. On the distribution of the two-sample cramer-von mises criterion. *The Annals of Mathematical Statistics*, 33(3):1148–1159, 1962.
- [6] Mayank Bansal, Alex Krizhevsky, and Abhijit Ogale. Chauffeurnet: Learning to drive by imitating the best and synthesizing the worst. *ArXiv*, 2018.
- [7] Daniel J. Benjamin, James O. Berger, Magnus Johannesson, Brian A. Nosek, E-J Wagenmakers, Richard Berk, Kenneth A. Bollen, Björn Brembs, Lawrence Brown, Colin Camerer, David Cesarini, Christopher D. Chambers, Merlise Clyde, Thomas D. Cook, Paul de Boeck, Zoltan Dienes, Anna Dreber, Kenny Easwaran, Charles Effer-son, Ernst Fehr, Fiona Fidler, Andy P. Field, Malcolm Forster, Edward I. George, Richard Gonzalez, Steven Goodman, Edwin Green, Donald P. Green, Anthony G. Greenwald, Jarrod D. Hadfield, Larry V. Hedges, Leonhard Held, Teck Hua Ho, Herbert Hoijtink, Daniel J. Hruschka, Kosuke Imai, Guido Imbens, John P. A. Ioan-nidis, Minjeong Jeon, James Holland Jones, Michael Kirchler, David Laibson, John List, Roderick Little, Arthur Lupia, Edouard Machery, Scott E. Maxwell, Michael McCarthy, Don A. Moore, Stephen L. Morgan, Marcus Munafó, Shinichi Nakagawa, Brendan Nyhan, Timothy H. Parker, Luis Pericchi, Marco Perugini, Jeff Rouder, Ju-dith Rousseau, Victoria Savalei, Felix D. Schönbrodt, Thomas Sellke, Betsy Sinclair, Dustin Tingley, Trisha van Zandt, Simine Vazire, Duncan J. Watts, Christopher Win-ship, Robert L. Wolpert, Yu Xie, Cristobal Young, Jonathan Zinman, and Valen E. Johnson. Redefine statistical significance. *Nature human behaviour*, 2(1):6–10, 2018.
- [8] Mario Berk, Olaf Schubert, Hans-Martin Kroll, Boris Buschardt, and Daniel Straub. Assessing the safety of environment perception in automated driving vehicles. *SAE International Journal of Transportation Safety*, 8(1):49–74, 2020.
- [9] Holger Caesar, Varun Bankiti, Alex H. Lang, Sourabh Vora, Venice Erin Liong, Qiang Xu, Anush Krishnan, Yu Pan, Giancarlo Baldan, and Oscar Beijbom. nuscenes: A multimodal dataset for autonomous driving. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020.
- [10] Holger Caesar, Juraj Kabzan, Kok Seang Tan, Whye Kit Fong, Eric M. Wolff, Alex Lang, Luke Fletcher, Oscar Beijbom, and Sammy Omari. nuplan: A closed-loop ml-based planning benchmark for autonomous vehicles. *ArXiv*, 2021.

- [11] Bill Canis. Autonomous vehicles: Emerging policy issues. *Congressional Research Service*, 23.05.2017.
- [12] K. Cho, T. Ha, G. Lee, and S. Oh. Deep predictive autonomous driving using multi-agent joint trajectory prediction and traffic rules. In *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 2076–2081, 2019.
- [13] VVM Consortium. Verification & validation methods project. <https://www.vvm-projekt.de/>, last accessed 17.07.2023.
- [14] Boyang Deng, C. Qi, Mahyar Najibi, Thomas A. Funkhouser, Yin Zhou, and Drago Anguelov. Revisiting 3d object detection from an egocentric perspective. In *NeurIPS*, 2021.
- [15] Nachiket Deo, Eric Wolff, and Oscar Beijbom. Multimodal trajectory prediction conditioned on lane-graph traversals. In Aleksandra Faust, David Hsu, and Gerhard Neumann, editors, *Proceedings of the 5th Conference on Robot Learning*, volume 164 of *Proceedings of Machine Learning Research*, pages 203–212. PMLR, 2022.
- [16] Alexey Dosovitskiy, German Ros, Felipe Codevilla, Antonio Lopez, and Vladlen Koltun. Carla: An open urban driving simulator. *ArXiv*, 2017.
- [17] Scott M. Ettinger, Shuyang Cheng, Benjamin Caine, Chenxi Liu, Hang Zhao, Sabeek Pradhan, Yuning Chai, Benjamin Sapp, C. Qi, Yin Zhou, Zoey Yang, Aurelien Chouard, Pei Sun, Jiquan Ngiam, Vijay Vasudevan, Alexander McCauley, Jonathon Shlens, and Drago Anguelov. Large scale interactive motion forecasting for autonomous driving : The waymo open motion dataset. *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 9690–9699, 2021.
- [18] Andreas Geiger, Philip Lenz, and Raquel Urtasun. Are we ready for autonomous driving? the kitti vision benchmark suite. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3354–3361, 2012.
- [19] Sven Hallerbach, Yiqun Xia, Ulrich Eberle, and Frank Koester. Simulation-based identification of critical scenarios for cooperative and automated vehicles. *SAE International Journal of Connected and Automated Vehicles*, 1(2):93–106, 2018.
- [20] Franziska Henze, Dennis Fabender, and Christoph Stiller. Identifying admissible uncertainty bounds for the input of planning algorithms. *IEEE Transactions on Intelligent Vehicles*, page 1, 2021.
- [21] Michael Hoss, Maike Scholtes, and Lutz Eckstein. A review of testing object-based environment perception for safe automated driving. *Automotive Innovation*, 5(3):223–250, 2022.
- [22] Sebastian Houben, Stephanie Abrecht, Maram Akila, Andreas Bär, Felix Brockherde, Patrick Feifel, Tim Fingscheidt, Sujun Sai Gannamaneni, Seyed Eghbal Ghobadi, Ahmed Hammam, Anselm Haselhoff, Felix Hauser, Christian Heinzemann, Marco Hoffmann, Nikhil Kapoor, Falk Kappel, Marvin Klingner, Jan Kronenberger, Fabian

- Küppers, Jonas Löhdefink, Michael Mlynarski, Michael Mock, Firas Mualla, Svetlana Pavlitskaya, Maximilian Poretschkin, Alexander Pohl, Varun Ravi-Kumar, Julia Rosenzweig, Matthias Rottmann, Stefan Rüping, Timo Sämann, Jan David Schneider, Elena Schulz, Gesina Schwalbe, Joachim Sicking, Toshika Srivastava, Serin Varghese, Michael Weber, Sebastian Wirkert, Tim Wirtz, and Matthias Woehrle. Inspect, understand, overcome: A survey of practical methods for ai safety. In Tim Fingscheidt, Hanno Gottschalk, and Sebastian Houben, editors, *Deep Neural Networks and Data for Automated Driving*, pages 3–78. Springer Cham, 2022.
- [23] John Houston, Guido Zuidhof, Luca Bergamini, Yawei Ye, Long Chen, Ashesh Jain, Sammy Omari, Vladimir Iglovikov, and Peter Ondruska. One thousand and one hours: Self-driving motion prediction dataset. *ArXiv*, 2020.
- [24] Saurabh Jha, Shengkun Cui, Zbigniew Kalbarczyk, and Ravishankar K. Iyer. Watch out for the safety-threatening actors: Proactively mitigating safety hazards. *ArXiv*, 2022.
- [25] Philipp Junietz, Udo Steininger, and Hermann Winner. Macroscopic safety requirements for highly automated driving. *Transportation Research Record*, 2673(3):1–10, 2019.
- [26] Peng Liu, Run Yang, and Zhigang Xu. How safe is safe enough for self-driving vehicles? *Risk analysis : an official publication of the Society for Risk Analysis*, 39(2):315–325, 2019.
- [27] Maria Lyssenko, Christoph Gladisch, Christian Heinzemann, M. Woehrle, Rudolph Triebel, and R. Bosch. From evaluation to verification: Towards task-oriented relevance metrics for pedestrian detection in safety-critical domains. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, pages 38–45, 2021.
- [28] Jiageng Mao, Minzhe Niu, Chenhan Jiang, Hanxue Liang, Jingheng Chen, Xiaodan Liang, Yamin Li, Chaoqiang Ye, Wei Zhang, Zhenguo Li, Jie Yu, Hang Xu, and Chunjing Xu. One million scenes for autonomous driving: Once dataset. *ArXiv*, 2021.
- [29] Jiageng Mao, Shaoshuai Shi, Xiaogang Wang, and Hongsheng Li. 3d object detection for autonomous driving: A comprehensive survey. *International Journal of Computer Vision*, 2023.
- [30] Ken Mori, Kai Storms, and Steven Peters. Conservative estimation of perception relevance of dynamic objects for safe trajectories in automotive scenarios. *ArXiv*, 2023.
- [31] nachiket92. Pgp: Multimodal trajectory prediction conditioned on lane-graph traversals, 2022. <https://github.com/nachiket92/PGP>, last accessed 23.05.2023.
- [32] nuScenes. nusenes prediction task: Leaderboard, 2020. <https://www.nuscenes.org/prediction>, last accessed 23.05.2023.

- [33] PEGASUS Project. Pegasus method: An overview, 2019. <https://www.pegasusprojekt.de/files/tmpl/Pegasus-Abschlussveranstaltung/PEGASUS-Gesamtmethode.pdf>, last accessed 28.06.2022.
- [34] Jonah Philion and S. Fidler. Lift, splat, shoot: Encoding images from arbitrary camera rigs by implicitly unprojecting to 3d. In *ECCV*, 2020.
- [35] Jonah Philion, Amlan Kar, and S. Fidler. Learning to evaluate perception models using planner-centric metrics. *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 14052–14061, 2020.
- [36] Robin Philipp, Hedan Qian, Lukas Hartjen, Fabian Schuldt, and Falk Howar. Simulation-based elicitation of accuracy requirements for the environmental perception of autonomous vehicles. In Tiziana Margaria and Bernhard Steffen, editors, *Leveraging Applications of Formal Methods, Verification and Validation*, pages 129–145, Cham, 2021. Springer International Publishing.
- [37] Robin Philipp, Jana Rehbein, Felix Grun, Lukas Hartjen, Zhijing Zhu, Fabian Schuldt, and Falk Howar. Systematization of relevant road users for the evaluation of autonomous vehicle perception. In *2022 IEEE International Systems Conference (SysCon)*, pages 1–8.
- [38] Robin Philipp, Fabian Schuldt, and Falk Howar. Functional decomposition of automated driving systems for the classification and evaluation of perceptual threats. In *13. Uni-DAS eV Workshop Fahrerassistenz und automatisiertes Fahren 2020*. Walting, 2020.
- [39] Charles R. Qi, Yin Zhou, Mahyar Najibi, Pei Sun, Khoa Vo, Boyang Deng, and Dragomir Anguelov. Offboard 3d object detection from point cloud sequences. In *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, pages 6130–6140, 2021.
- [40] A. Sadat, S. Casas, Mengye Ren, X. Wu, Pranaab Dhawan, and R. Urtasun. Perceive, predict, and plan: Safe motion planning through interpretable semantic representations. In *ECCV*, 2020.
- [41] Valerij Schöнемann, Mara Duschek, and Hermann Winner. Maneuver-based adaptive safety zone for infrastructure-supported automated valet parking. *2184-495X*, 2019.
- [42] Pei Sun, Henrik Kretzschmar, Xerxes Dotiwalla, Aurelien Chouard, Vijaysai Patnaik, Paul Tsui, James Guo, Yin Zhou, Yuning Chai, Benjamin Caine, Vijay Vasudevan, Wei Han, Jiquan Ngiam, Hang Zhao, Aleksei Timofeev, Scott Ettinger, Maxim Krivokon, Amy Gao, Aditya Joshi, Sheng Zhao, Shuyang Cheng, Yu Zhang, Jonathon Shlens, Zhifeng Chen, and Dragomir Anguelov. Scalability in perception for autonomous driving: Waymo open dataset. *ArXiv*, 2019.
- [43] Sever Topan, Karen Leung, Yuxiao Chen, Pritish Tupekar, Edward Schmerling, Jonas Nilsson, Michael Cox, and Marco Pavone. Interaction-dynamics-aware perception zones for obstacle detection safety evaluation. *2022 IEEE Intelligent Vehicles Symposium (IV)*, pages 1201–1210, 2022.

- [44] José Luis Vázquez, Alexander Liniger, Wilko Schwarting, Daniela Rus, and Luc van Gool. Deep interactive motion prediction and planning: Playing games with motion prediction models. *ArXiv*, 2022.
- [45] Georg Volk, Jörg Gamberdinger, Alexander von Betnuth, and O. Bringmann. A comprehensive safety metric to evaluate perception in autonomous systems. *2020 IEEE 23rd International Conference on Intelligent Transportation Systems (ITSC)*, pages 1–8, 2020.
- [46] Oliver Willers, Sebastian Sudholt, Shervin Raafatnia, and Stephanie Abrecht. Safety concerns and mitigation approaches regarding the use of deep learning in safety-critical perception tasks. In António Casimiro, editor, *Computer safety, reliability, and security*, volume 12235 of *Lecture Notes in Computer Science*, pages 336–350. Springer, Cham, 2020.
- [47] Benjamin Wilson, William Qi, Tanmay Agarwal, John Lambert, Jagjeet Singh, Siddhesh Khandelwal, Bowen Pan, Ratnesh Kumar, Andrew Hartnett, Jhony Kaese-model Pontes, Deva Ramanan, Peter Carr, and James Hays. Argoverse 2: Next generation datasets for self-driving perception and forecasting. *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks*, 1, 2021.
- [48] Mirja Wolf, Luiz R. Douat, and Michael Erz. Safety-aware metric for people detection. In *2021 IEEE International Intelligent Transportation Systems Conference (ITSC)*, pages 2759–2765. IEEE, 2021.

Operator Monitoring zur Absicherung von Teleoperation und automatisiertem Fahren: Ergebnisse aus dem DFG-Schwerpunktprogramm *CoInCar* sowie einer Disseration

Nicolas Herzberger¹, Marcel Usai² und Frank Flemisch³

Zusammenfassung: Trotz bedeutender technischer Fortschritte im Kontext des automatisierten Fahrens sind aktuelle Systeme noch weit davon entfernt, jede auftretende Situation selbstständig lösen zu können. Neben dem Lösungsweg des minimalrisikoreichen Manövers (MRM) stehen insbesondere Transitionen der Fahraufgabe an die Fahrperson (Takeover Request, TOR) oder an einen Teleoperator im Fokus aktueller Forschungsvorhaben. In beiden Fällen ist es jedoch notwendig einzuschätzen, ob der jeweilige Operator, im Fahrzeug oder einer entfernten Leitwarte, die Fahraufgabe sicher übernehmen kann. Der Beitrag konzentriert sich nach einer kurzen Einführung auf aktuelle Studienergebnisse aus dem grade abgeschlossenen DFG-Schwerpunktprogramm 1835 *CoInCar* und aus einer Dissertation zum Operator Monitoring, sowie auf deren Bedeutung für die Absicherung von Transitionen im Kontext des automatisierten Fahrens sowie der Teleoperation.

Schlüsselwörter: Operator Monitoring, automatisiertes Fahren, Teleoperation, Teledriving, Teleassistenz, kooperative Automation, Hochautomation

1 Einleitung

Seit dem ersten automatisierten Prototyp VaMoRs (Versuchsfahrzeug für autonome Mobilität und Rechnersehen), der 1992 in Deutschland mit beinahe 100 km/h auf einem gesperrten Autobahnabschnitt fuhr (Projekt PROMETHEUS [1]), über den vollständig selbstfahrenden VW Touareg *Stanley* der 2005 die DARPA Grand Challenge gewann [2] hin zu der ersten Zulassung eines SAE Level 3 Systems (Mercedes DRIVE PILOT [3]) Ende 2021, haben sich die technischen Möglichkeiten in den vergangenen Jahren stetig entwickelt. Parallel dazu haben aus der begleitenden Forschung viele bedeutende Grundlagen wie die Kooperation zwischen Mensch und Maschine [4], [5], [6], die

¹ Nicolas Herzberger ist Gruppenleiter am IAW der RWTH Aachen University sowie am Fraunhofer FKIE, Fraunhoferstr. 20, 53343 Wachtberg (e-mail: nicolas.herzberger@fkie.fraunhofer.de)

² Marcel Usai ist wissenschaftlicher Mitarbeiter an der RWTH Aachen University sowie am Fraunhofer FKIE, Fraunhoferstr. 20, 53343 Wachtberg (e-mail: marcel.usai@fkie.fraunhofer.de)

³ Frank Flemisch ist Professor am IAW der RWTH Aachen University sowie Abteilungsleiter am Fraunhofer FKIE, Fraunhoferstr. 20, 53343 Wachtberg (e-mail: frank.flemisch@fkie.fraunhofer.de)

So sind aktuelle Systeme bislang nicht in der Lage, alle Situationen selbstständig lösen zu können und werden dies aller Voraussicht nach auch in naher Zukunft nicht können. Als Beispiel für eine „unlösbare“ Situation soll hier ein unzulässig in zweiter Reihe geparkter Transporter dienen, bei gleichzeitig durchgezogenen Fahrstreifenmittellinie. Das Überholen über die Gegenfahrbahn wäre einem automatisierten System in dieser Situation entsprechend der Straßenverkehrsordnung nicht gestattet und es würde auf unbestimmte Zeit hinter dem Lieferwagen warten. Hier müsste der Passagier dann wieder zur Fahrperson werden, die Fahraufgabe übernehmen, das Hindernis umfahren und die Fahraufgabe wieder an die Automation abgeben. Alternativ, besonders auch im Kontext des ÖPNV wie z.B. bei Shuttles/Taxen etc. werden jedoch auch Transitionen der Fahraufgabe an einen Teleoperator erforscht. In beiden Fällen ist es jedoch aus sicherheitstechnischen Gründen absolut essentiell einschätzen zu können, ob die Fahrperson (lokal im Fahrzeug oder ein externer Teleoperator) in der Lage ist, eine sichere Übernahme leisten zu können.



Wie im linken Teil der Abbildung dargestellt, werden die Fähigkeiten der Fahrperson (orange) kontinuierlich mit denen des automatisierten Fahrsystems (blau) abgeglichen. Daraus ergibt sich entweder eine Überlappung, eine Sicherheitsreserve (wie rechts in der Abbildung dargestellt), oder ein Delta, d.h. ein kritisches Kontrolldefizit. Zur Einschätzung der Fähigkeiten der Fahrperson bzw. des Teleoperators werden im folgenden Kapitel aktuelle Studienergebnisse aus dem DFG SPP 1835 *CoInCar* sowie einer Dissertation [14] zum Operator Monitoring vorgestellt.

2 Operator Monitoring

Bereits vor der Einführung (teil-)automatisierter Fahrfunktionen wurden erste Systeme zur Unterstützung und Überwachung des Operators in Fahrzeugen eingesetzt. So starteten beispielsweise Müdigkeitserkennung zunächst rein zeitabhängig und entwickelten sich mit der Zeit über einfache Systeme, die ruckartiger werdendes Lenkverhalten detektierten, und Spurhalte-Erkennungssysteme, die Abweichungen von der Nulllinie registrierten, hin zu komplexen Systemen, die Veränderungen im Lenkverhalten zu individuellen Benchmarks mit kamerabasierten Informationen zum Außengeschehen und dem Blick der Fahrperson koppeln. Aber besonders mit zunehmender Automatisierung nahm die Forschung zum Thema Operator Monitoring deutlich zu: Von diversen physiologischen Parametern (EEG, EKG, EDA etc.) über Blickrichtungsmessung und Pupillometrie hin zu kombinierten Maßen, wie z.B. Ansätzen zur Emotionsdetektion finden sich unterschiedlichste Konzepte und Prototypen in der Entwicklung oder bereits im Markt. Eine fundierte Einschätzung dieser Ansätze ermöglicht beispielsweise die Metastudie [15], welche auf Blickrichtungsmessungen basierende Systeme für die Technologie mit dem größten Potential hält.

Darauf aufbauend wurde im Rahmen der Dissertation [14] zunächst ein erster Schätzer entwickelt, welcher Blicke auf die Verkehrssituation zu Grunde legte. Dieser Ausgangspunkt wurde gewählt, da der Blick auf die Straße zumindest ein notwendiges Kriterium für den Aufbau des Situationsbewusstseins darstellt, wenngleich dies kein hinreichendes Kriterium darstellen muss (*looking but not seeing*).

Ausgehend von dem Ziel, diesen grundlegenden Involvierungsschätzer um zusätzliche Kriterien zu erweitern, zeigte sich, dass es bislang keine, zumindest keine öffentlich zugänglichen Ergebnisse gibt, welche Kriterien erfasst werden müssen, um eine (fehlende) Übernahmefähigkeit valide zu erfassen.

Um diese Kriterien zu identifizieren, wurde zwei Studien (statischer Fahrsimulator $N = 40$ [15] und Realverkehr $N = 20$ [16]) durchgeführt, die einen alternativen Ansatz zur Beurteilung der Involvierung von Fahrpersonen verfolgten. Dieser Ansatz beruhte auf der Annahme, dass Menschen auf dem Beifahrersitz in der Lage sind den Zustand und die aktuelle Aufmerksamkeit von Fahrpersonen zuverlässig einschätzen zu können – vergleichbar Fahrlehrenden, die die Aufmerksamkeit und Fähigkeiten ihrer Fahrlehrenden einschätzen und basierend darauf entscheiden, ob und wann es nötig ist in die Fahraufgabe einzugreifen bzw. wann die Verkehrssituation von diesen selbstständig beherrscht werden kann [16].

Zur Identifikation der Bewertungskriterien wurden im Rahmen der Studien Fahrpersonen vor einer Übernahmeaufforderung aus zwei Perspektiven (1. Sicht auf die Umgebungssituation und 2. die Fahrperson) gefilmt, siehe Abbildung 2. Diese Videoaufzeichnungen wurden anschließend Teilnehmenden in einer Onlinebefragung (statischer Fahrsimulator $N = 80$ [15] und Realverkehr $N = 233$ [16]) präsentiert zusammen mit den Fragen, ob die Fahrperson im Video die Fahraufgabe sicher übernehmen könnte und woran diese Entscheidung festgemacht wurde. Diese Einschätzungen wurden anschließend mit den tatsächlichen Übernahmeleistungen abgeglichen.



Abbildung 2: Zusammengefügte synchronisierte Videoausschnitte aus zwei Perspektiven [15].
Links: Fahrzeuginnenraum und Situation vor dem Fahrzeug. Rechts: Frontansicht der Fahrperson.

Es zeigte sich, dass die Teilnehmenden an den Onlinestudien keine valide Bewertung und somit keinen nutzbaren Kriterienpool liefern konnten. Dies lag vor allem daran, dass die Einschätzungen besonders auf dem Kriterium *Blick auf die Straße* beruhten. Besonders die Bewertung der Fahrpersonen im Fahrsimulator offenbarte hierbei aber eine besondere Herausforderung. So wurden den Fahrpersonen unterschiedliche fahrfremde Tätigkeiten angeboten, wie ein Tetris-Spiel oder ein Skype-Telefonat, in dem Quizfragen beantwortet werden sollten, um ein anstrengendes Telefonat zu simulieren.

Grade die zweite Aufgabe, das Telefonat, führte jedoch dazu, dass die Teilnehmenden zwar auf die Straße blickten, aber so stark abgelenkt waren, dass keine sichere Übernahme durchgeführt werden konnte (siehe Abbildung 2, rechts).

Da der gewählte Ansatz, die Beobachtung durch Beifahrende, nicht die gewünschten Bewertungskriterien lieferte und auch kein alternativer veröffentlichter Kriterienkatalog zur Verfügung stand, wurde mit dem vorrangigen Ziel der Erfassung der (fehlenden) Übernahmefähigkeit, daraufhin die Methode des Diagnostischen Takeover Requests [17] entwickelt.

Dabei erfolgt die Erfassung fehlender Übernahmefähigkeit mittels Klassifikationen von Orientierungsreaktionen auf eine Übernahmeaufforderung (Takeover Request, TOR), welche zunächst bei einer Vielzahl an Fahrpersonen, zusammen mit der anschließenden Übernahmequalität aufgezeichnet und ausgewertet werden. Basierend auf diesem Datensatz lassen sich post-hoc sichere und somit gute, von risikoreicheren, schlechten Übernahmen unterscheiden. Nachdem diese Aufteilung erfolgt ist, können die zuvor

gezeigten Orientierungsreaktionen auf den TOR analysiert werden, mit dem Ziel, trennscharfe Orientierungsreaktionen vor sicheren und unsicheren Übernahmen zu identifizieren [18], [19].

Zur Überprüfung der Methode wurden zwei Fahrsimulatorstudien (Studie 1: $N = 50$, statischer Fahrsimulator des Ika der RWTH Aachen University, Studie 2: $N = 38$, statischer Fahrsimulator des IAW der RWTH Aachen University, siehe Abbildung 3) durchgeführt.



Abbildung 3: Statischer Fahrsimulator am IAW der RWTH Aachen University.

In beiden Studien durchfuhren unterschiedliche Teilnehmende das identische Szenario, siehe Abbildung 1 rechts, in welchem sie nach einer automatisierten Fahrt (15 min, 50% der Teilnehmenden fuhren Tetis als fahrfremde Tätigkeit aus) einen TOR erhielten. Anschließend wurden die Blickabfolgen in die möglichen Areas of Interest (AOI) entsprechend der ISO 15007:2020 sowie die jeweilige Übernahmequalität ausgewertet. Dabei konnten mehrere trennscharfe Blickfolgen, auch über beide Stichproben hinweg, identifiziert werden. Die Methode des Diagnostischen Takeover Requests scheint damit ein erster vielversprechender Ansatz zur Erfassung kritischer fehlender Übernahmefähigkeit zu sein.

3 Implikationen für Transitionen an die Fahrperson oder einen Teleoperator

Die Ergebnisse der Studien skizzieren mit dem Diagnostischen Takeover Request einen vielversprechenden Ansatz für die Absicherung der Transitionen an die Fahrperson oder einen Teleoperator. Hinsichtlich der Transitionen und der zugehörigen Absicherungsanforderungen müssen diese jedoch differenziert betrachtet werden:

So werden Transitionen im Normalbetrieb und an unkritischen Systemgrenzen, wie z.B. eine bevorstehende Autobahnausfahrt, wenn die Operation Desing Domain (ODD) nur die Autobahn beinhaltet, sehr viel weniger zeit- und sicherheitskritisch sein als Systemausfälle oder plötzlich auftretende Systemgrenzen. Und auch für die Teleoperation müssen, die Anforderungen an Operator Monitoring Systeme betreffend, unterschiedlich kritische Anwendungsfälle unterschieden werden: So werden im Bereich des niederautomatisierten Teledrivings, in welchem die Fahrzeuge dauerhaft aus der Ferne gesteuert werden, kaum bis gar keine Transitionen stattfinden. Hier besteht die Sinnhaftigkeit für Operator Monitoring Systeme eher im Bereich der Müdigkeits-/ Aufmerksamkeitsüberwachung und der Lenk- bzw. Arbeitszeiten. Deutlich (zeit-) kritischer hingegen ist der Bereich des teil- und hochautomatisierten Teledrivings, sowie der Teleassistenz, in welcher ein Teleoperator nur „on demand“ hinzugeschaltet wird. Dieser kann dann entweder das Fahrzeug direkt steuern, was hohe Anforderungen an beispielsweise Latenzen mit sich bringt, oder ausschließlich manöverbasiert interagieren, indem einzelne Fahrmanöver, z.B. ein Überholvorgang lediglich freigegeben werden, die eigentliche Steuerung aber weiterhin lokal von der Fahrzeugautomation übernommen wird [20]. Die Transitionen zu einem Teleassistent sind daher eher mit denen an eine lokale Fahrperson vergleichbar. Dennoch eröffnet das Operator-Monitoring im Bereich der Teleoperation eine Vielzahl neuer Fragestellungen, wie z.B. In welchem Kontext und in welchem Umfang muss, auch aus zulassungsrechtlichen Gründen, eine Überwachung des Teleoperators erfolgen? Was müssen diese Systeme erfassen? Ist ein verringertes Situationsbewusstsein des Teleoperators z.B. auf Grund der räumlichen Entkopplung zu erwarten? Geht dies mit Reaktionsverzögerungen einher? Lassen sich bisherige Operator-Monitoring Konzepte und Methoden übertragen?

Grade die Übertragbarkeit bestehender Ansätze erscheint jedoch fraglich. So lassen sich z.B. die bestehenden lenkungsbasierten Müdigkeitserfassungssysteme kaum einsetzen, da durch Latenzschwankungen von einer hohen Varianz im Lenkverhalten auszugehen ist. Und auch die Orientierungsmuster, wie sie z.B. während des Diagnostischen TORs analysiert werden, lassen sich nicht ohne weiteres übertragen, da sich die (zusätzlichen) Ansichten des Teleoperators auch in der Anordnung von denen der Fahrperson im Fahrzeug unterscheiden.

Die offenen Fragestellungen sollten jedoch in naher Zukunft adressiert werden, da bereits heute erste Fahrzeuge mit Ausnahmegenehmigungen im Straßenverkehr unterwegs sind und erste Unfälle, die durch ein effektives Operator Monitoring hätten vermieden werden können, leicht das Vertrauen in die Teleoperation als zukünftigen Use-Case und als Geschäftsmodell verspielen könnten.

4 Fazit und Ausblick

Ausgangspunkt für den vorliegenden Beitrag waren große technische Fortschritte des automatisierten Fahrens sowie im Bereich der Teleoperation, die gleichzeitig wieder neue Herausforderungen für die sichere Mensch-Maschine-Kooperation und -Interaktion bedeuten. Die Ergebnisse aus den Studien zeigen, dass auch bei geringen Übernahmezeiten erste Abschätzungen zur Absicherung der Transition getroffen werden

können. Jedoch gilt es dabei zu berücksichtigen, dass nur eine Situation im Detail getestet wurde und auf die einzelnen Orientierungsmuster auf Grund der hohen Anzahl an möglichen Kombinationen nur geringe Fallzahlen der Stichprobe entfallen. Dennoch bietet die Methode des Diagnostischen Takeover Requests in Verbindung mit den Confidence Horizons gute Möglichkeiten, die Orientierungsmuster weiter zu verdichten: So beschäftigt sich eine anschließende Dissertation aktuell mit der Ergänzung um weitere Bestandteile der Orientierung, wie dem Kontakt zu Stellteilen, oder der Verlagerung der Körperhaltung. Erste Ergebnisse deuten auch hier auf trennscharfe Muster hin, die zur weiteren Absicherung der Transitionen genutzt werden können. Inwiefern sich diese Orientierungsmuster auch auf Übergaben an einen Teleoperator übertragen lassen, müssen zukünftige Studien adressieren. Die Teleoperation gerade in Verbindung mit Automation ist aktuell ein Themenfeld mit einer hohen Dynamik, was begleitende absichernde Forschungsvorhaben gerade auch zu Transitionen unabdingbar macht.

Literatur

- [1] Lossau, N. (2017). Das erste autonome Auto kostete 200.000 D-Mark. Welt. <https://www.welt.de/wissenschaft/article169604489/Das-erste-autonome-Auto-kostete-200-000-D-Mark.html>
- [2] Thrun, S., Montemerlo, M., Dahlkamp, H., Stavens, D., Aron, A., Diebel, J., Fong, P., Gale, J., Halpenny, M., Hoffmann, G., Lau, K., Oakley, C., Palatucci, M., Pratt, V., Stang, P., Strohband, S., Dupont, C., Jendrossek, L.-E., Koelen, C., Markey, C., Rummel, C., van Nieuwerkerk, J., Jensen, E., Alessandrini, P., Bradski, G., Davies, B., Ettinger, S., Kaehler, A., Nefian, A. & Mahoney, P. (2006). Stanley: The robot that won the DARPA Grand Challenge. *Journal of Field Robotics*, 23(9), 661–692.
- [3] Hubik, F. (9. Dezember 2021). Daimler erhält Level-3-Freigabe: S-Klasse fährt teils autonom. Handelsblatt. <https://www.handelsblatt.com/unternehmen/industrie/autonomes-fahren-daimler-erhaelt-level-3-freigabe-s-klasse-darf-autonom-auf-autobahnen-fahren/27877760.html?ticket=ST-8718105-AbBgv3bH0eO9SspLduDG-cas01.example.org>
- [4] Flemisch, F. O., Adams, C. A., Conway, S. R., Goodrich, K. H., Palmer, M. T. & Schutte, P. (2003). The H-Metaphor as a guideline for vehicle automation and interaction.
- [5] Hoc, J.-M. & Pacaux-Lemoine, M.-P. (1998). Cognitive Evaluation of Human-Human and Human-Machine Cooperation Modes in Air Traffic Control. *The International Journal of Aviation Psychology*, 8(1), 1–32.
- [6] Pacaux-Lemoine, M.-P. & Flemisch, F. O. (2016). Layers of Shared and Cooperative Control, assistance and automation. *IFAC-PapersOnLine*, 49(19), 159–164.
- [7] Flemisch, F. O., Bengler, K., Bubb, H., Winner, H. & Bruder, R. (2014). Towards cooperative guidance and control of highly automated vehicles: H-Mode and Conduct-by-Wire. *Ergonomics*, 57(3), 343–360.
- [8] Schneemann, F. & Diederichs, F. (2019). Action prediction with the Jordan model of human intention: a contribution to cooperative control. *Cognition, Technology & Work*, 21(4), 711–721
- [9] Gasser, T. M., Arzt, C., Ayoubi, M., Bartels, A., Bürkle, L., Eier, J., Flemisch, F., ... & Vogt, W. (2012). Rechtsfolgen zunehmender Fahrzeugautomatisierung: Gemeinsamer Schlussbericht der Projektgruppe ; [Bericht zum Forschungsprojekt F 1100.5409013.01. Berichte der Bundesanstalt für Strassenwesen F, Fahrzeugtechnik: Bd. 83. Wirtschaftsverl. NW Verl. für neue Wissenschaft.
- [10] SAE International (2021). Taxonomy and Definitions for Terms Related to Driving Automation Systems for On-Road Motor Vehicles (No. J3016_202104). 400 Commonwealth Drive, Warrendale, PA, United States. SAE International.
- [11] Herzberger, N. D., Schwalm, M., Reske, M., Woopen, T., & Eckstein, L. (2019). *Mobilitätskonzepte der Zukunft- Ergebnisse einer Befragung von 619 Personen in Deutschland im Rahmen des Projekts UNICARagil*. Universitätsbibliothek der RWTH Aachen.
- [12] Flemisch, F., Schwalm, M. (2019). Systemergonomie für kooperativ interagierende Automobile: Nachvollziehbarkeit des Automationsverhaltens im Normalbetrieb, an Systemgrenzen und bei Systemausfall. DFG Folgeantrag.
- [13] Herzberger, N., Usai, M. & Flemisch, F. (2022). Confidence Horizon for a Dynamic Balance between Drivers and Vehicle Automation: First Sketch and Application. In AHFE International, Human Factors in Transportation. AHFE International.

- [14] Herzberger, N.D. (2023). Erfassung der Übernahmefähigkeit von Fahrpersonen im Kontext des automatisierten Fahrens. Dissertation. Shaker Verlag.
- [15] Herzberger, N. D., Schwalm, M., Flemisch, F. O., Schmidt, E. & Sitter, A. (2019). Erfassung der Fahrerübernahmefähigkeit im automatisierten Fahren anhand von Fahrerbeobachtungen. In *Mensch-Maschine-Mobilität 2019* (S. 53–66). VDI Verlag. <https://doi.org/10.51202/9783181023600-53>
- [16] Meyer, M.-L. (2017). *Fahrerzustand - Der Beifahrer als Maßstab*. Unveröffentlichte Masterarbeit. RWTH Aachen University.
- [17] Schories, L., Erggelet, M., Schwalm, M. & Herzberger, N. D. (2018). *Verfahren zum Feststellen einer Übernahmefähigkeit eines Fahrzeugnutzer eines Fahrzeugs* (DE 10 2018 007 508 A8). Deutschland.
- [18] Herzberger, N. D., Eckstein, L. & Schwalm, M. (2018). Detection of Missing Takeover Capability by the Orientation Reaction to a Takeover Request. In *Aachener Kolloquium Fahrzeug- und Motorentechnik GbR* (Hrsg.), 27. Aachen Colloquium Automobile and Engine Technology (S. 1231–1240).
- [19] Schwalm, M., & Herzberger, N. D. (2018). Die Erfassung des Fahrerzustands als Voraussetzung für höher automatisierte Fahrfunktionen–Eine kritische Diskussion und ein Lösungsvorschlag. In 12. Workshop Fahrerassistenzsysteme. Walting (Albmühltal) (Vol. 26, No. 28.09, p. 2018)
- [20] Herzberger, N., Wasser, J., Flemisch, F. (2022). Control Centers for Maneuver-based Teleoperation of Highly Automated and Autonomous Vehicles: System Model and Requirements. In: Katie Plant and Gesa Praetorius (eds) *Human Factors in Transportation*. AHFE (2022) International Conference. AHFE Open Access, vol 60. AHFE International, USA.

Uni-DAS

Herausgeber

Klaus Bengler

Klaus Dietmayer

Lutz Eckstein

Meike Jipp

Markus Maurer

Steven Peters

Christoph Stiller

ISBN: 978-3-941543-74-4