

Scenario Parameter Generation Method and Scenario Representativeness Metric for Scenario-Based Assessment of Automated Vehicles

de Gelder, Erwin; Hof, Jasper; Cator, Eric; Paardekooper, Jan Pieter; Camp, Olaf Op den; Ploeg, Jeroen; de Schutter, Bart

DOI

[10.1109/TITS.2022.3154774](https://doi.org/10.1109/TITS.2022.3154774)

Publication date

2022

Document Version

Accepted author manuscript

Published in

IEEE Transactions on Intelligent Transportation Systems

Citation (APA)

de Gelder, E., Hof, J., Cator, E., Paardekooper, J. P., Camp, O. O. D., Ploeg, J., & de Schutter, B. (2022). Scenario Parameter Generation Method and Scenario Representativeness Metric for Scenario-Based Assessment of Automated Vehicles. *IEEE Transactions on Intelligent Transportation Systems*, 23(10), 18794-18807. <https://doi.org/10.1109/TITS.2022.3154774>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Scenario Parameter Generation Method and Scenario Representativeness Metric for Scenario-Based Assessment of Automated Vehicles

Erwin de Gelder¹, Jasper Hof², Eric Cator³, Jan-Pieter Paardekooper⁴, Olaf Op den Camp⁵, Jeroen Ploeg⁶, and Bart de Schutter⁷, *Fellow, IEEE*

Abstract—The development of assessment methods for the performance of Automated Vehicles (AVs) is essential to enable the deployment of automated driving technologies, due to the complex operational domain of AVs. One candidate is scenario-based assessment, in which test cases are derived from real-world road traffic scenarios obtained from driving data. Because of the high variety of the possible scenarios, using only observed scenarios for the assessment is not sufficient. Therefore, methods for generating additional scenarios are necessary. Our contribution is twofold. First, we propose a method to determine the parameters that describe the scenarios to a sufficient degree while relying less on strong assumptions on the parameters that characterize the scenarios. By estimating the probability density function (pdf) of these parameters, realistic parameter values can be generated. Second, we present the Scenario Representativeness (SR) metric based on the Wasserstein distance, which quantifies to what extent the scenarios with the generated parameter values are representative of real-world scenarios while covering the actual variety found in the real-world scenarios. A comparison of our proposed method with methods relying on assumptions of the scenario parameterization and pdf estimation shows that the proposed method can automatically determine the optimal scenario parameterization and pdf estimation. Furthermore, it is demonstrated that our SR metric can be used to choose the (number of) parameters that best describe a scenario. The presented method is promising, because the parameterization and pdf estimation can directly be applied to already available importance sampling strategies for accelerating the evaluation of AVs.

Index Terms—Intelligent vehicles, autonomous vehicles, vehicle safety, performance evaluation, automatic testing.

I. INTRODUCTION

AN ESSENTIAL aspect in the development of Automated Vehicles (AVs) is the assessment of the quality and performance of AV behavior with respect to safety, comfort, and efficiency [1]–[3]. Because public road tests are expensive and time consuming [4], [5], a scenario-based approach has been proposed [2], [6]–[11]. With a scenario-based approach, the response of the system-under-test is assessed in many scenarios and for the variations of these scenarios that occur in the real world. Here, a scenario describes the situation the system-under-test is in and how this situation develops over time (in Section III-A, a precise definition of the term scenario is provided). One of the advantages of a scenario-based approach is that the assessment can focus on the more challenging situations by selecting scenarios that are challenging for the system-under-test. As a source of information for the assessment scenarios, real-world driving data has been proposed, thereby guaranteeing that the scenarios represent real-world driving conditions [7]–[9].

For the scenario-based assessment approach, it is important that the generated scenarios are representative of scenarios that could happen in real life. In other words, the scenarios should be a representation of the real world [6]. Only then, the results of the assessment can be generalized to the performance of the system-under-test when operating in real life [10]. Furthermore, it is essential that the generated scenarios cover the same variety that is found in real life. Riedmaier *et al.* [6] argue that since an infinite number of situations occur in the real world, the scenario generation methods must provide a large number of variations in order to cover this infinite number of situations.

Our data-driven approach uses observed scenarios to generate parameter values that describe new scenarios. Instead of relying on a predetermined functional form of the signals, such as a vehicle's speed, and fitting parameters to this functional form, we employ a Singular Value Decomposition (SVD) [12] to determine in a data-driven manner the parameters that best describe the scenarios. Next,

Manuscript received June 30, 2021; revised November 10, 2021 and February 10, 2022; accepted February 22, 2022. The Associate Editor for this article was S.-H. Kong. (Corresponding author: Erwin de Gelder.)

Erwin de Gelder is with the Integrated Vehicle Safety Department, TNO, 5708 JZ Helmond, The Netherlands, and also with the Delft Center for Systems and Control, Delft University of Technology, 2628 CN Delft, The Netherlands (e-mail: erwin.degelder@tno.nl).

Jasper Hof is with the Department for Health Evidence, Radboud University Medical Center, 6525 GA Nijmegen, The Netherlands.

Eric Cator is with the Applied Stochastics Research Group, Radboud University, 6525 AJ Nijmegen, The Netherlands.

Jan-Pieter Paardekooper is with the Integrated Vehicle Safety Department, TNO, 5708 JZ Helmond, The Netherlands, and also with the Donders Institute for Brain, Cognition and Behaviour, Radboud University, 6500 GL Nijmegen, The Netherlands.

Olaf Op den Camp is with the Integrated Vehicle Safety Department, TNO, 5708 JZ Helmond, The Netherlands.

Jeroen Ploeg is with the Dynamics and Control Group, Department of Mechanical Engineering, Eindhoven University of Technology, 5600 MB Eindhoven, The Netherlands.

Bart de Schutter is with the Delft Center for Systems and Control, Delft University of Technology, 2628 CN Delft, The Netherlands.

Digital Object Identifier 10.1109/TITS.2022.3154774

the probability density function (pdf) of the parameters is estimated, such that the pdf can be used to sample the parameters to generate similar scenarios. To not assume a particular shape of the pdf, Kernel Density Estimation (KDE) [13], [14] is used for the pdf estimation. Furthermore, with KDE, the correlations that might exist among the parameters is modeled. This work also proposes a novel metric called the Scenario Representativeness (SR) metric for quantifying to what extent the generated scenarios are representative and cover the actual variety of real-world scenarios. More specifically, this metric uses the Wasserstein distance [15] to compare a set of generated scenarios with a set of observed scenarios.

This article is organized as follows. Section II reviews related works. In Section III, the approach for generating scenarios for the assessment of AVs is explained. Next, Section IV presents a novel metric for quantifying the performance of the scenario-generation method. A case study is performed in Section V. Section VI discusses relevant implications of our approach and some directions for future research. Conclusions of the paper are provided in Section VII.

II. RELATED WORKS

In this section, first, works concerning the generation of scenarios for the assessment of AVs are reviewed. Next, works related to the SR metric are reviewed.

A. Scenario Generation

The approaches to determine scenarios for the assessment of AVs can be categorized into three kinds [16]: scenarios based on observations of real-world traffic, scenarios based on the functionality that is being assessed, and a combination of these two approaches. The current paper focuses on the first approach.

In the literature, several methods are proposed to generate scenarios for the assessment based on real-world driving data. Lages *et al.* [17] proposed a method to construct scenarios in a virtual simulation environment by reconstructing the real-world scenarios observed by laser scanners. Zofka *et al.* [18] presented how recorded sensor data can be exploited to create scenarios that might lead to critical situations by modifying parameters of the recorded parameterized scenarios. Stepien *et al.* [19] generate scenarios by sampling scenario parameter values from generalized extreme value distributions, where the distribution parameters are fitted using scenario parameter values extracted from safety-critical scenarios observed in naturalistic driving data. In [10], [20]–[24], also parameterized scenarios were generated and, in addition, importance sampling techniques were presented that automatically generate scenarios in which the system-under-test shows (safety-)critical behavior. Other approaches to generate scenarios in which the system-under-test shows (safety-)critical behavior are Monte Carlo tree search [25] and genetic programming [26]. Schuldt *et al.* [27] provided a method to generate scenarios using combinatorial algorithms that should ensure that the test cases cover the variety of the possible situations the system-under-test could encounter

in real life. More recently, Spooner *et al.* [28] presented a Generative Adversarial Network (GAN) to generate pedestrian crossing scenarios.

In the existing literature, the scenario generation methods for the assessment of AVs have either one or more of the following shortcomings:

- Observed scenarios are replayed without adding more variations [17]. In this case, the total variety of scenarios that is found in real life will not be covered unless unrealistic amounts of data are gathered.
- The scenarios are oversimplified. For example, a vehicle's speed profile follows a predetermined functional form [10], [19], [21].
- Assumptions regarding the scenario parameter distributions are made that potentially compromise the quality of the scenarios. For example, the parameters are assumed to originate from a Gaussian [29] or generalized extreme value [19] distribution, and/or it is assumed that (some of) the parameters are uncorrelated [24].
- Because no pdf of the scenario parameters is known [18], no evaluation can be made of the performance of the system once deployed on the road, because it is unknown how realistic and likely the scenarios are.

In Section III, a method is proposed that overcomes these shortcomings.

B. Scenario Representativeness Metric

The generated scenarios should represent scenarios that could happen in real life. Whereas different approaches exist in the literature regarding the generation of scenarios for the assessment of AVs, less is known about the comparison of the generated scenarios with real-life traffic. From the mentioned sources in Section II-A, only [24] compared their generated scenarios with the ground truth from naturalistic driving data. Feng *et al.* [24] compared the distributions of vehicle speeds and bumper-to-bumper distances between the constructed scenarios and the ground truth. To quantify the similarity between the distributions, the Hellinger distance [30] and the mean absolute error were used. The disadvantages of this approach are that:

- 1) the generated scenarios may still be substantially different even though the distributions of the vehicle speeds and bumper-to-bumper distances are similar, and
- 2) only the marginal distributions are considered while the correlation between the vehicle speeds and bumper-to-bumper distances might be completely different.

Whereas little is known about comparing the generated scenarios for the scenario-based assessment of AVs with ground truth data, many similarity metrics for comparing two pdfs are known [30]. Well-known metrics are the Minkowski metric [30], which is a generalized version of the Euclidean distance, the f -divergence, which is a generalized version of both the Kullback-Leibler divergence [31] and the Hellinger distance [30], and the Wasserstein metric [15]. For practical reasons, this work uses the Wasserstein metric. As is shown in Section IV-B, the Wasserstein distance can be estimated using empirical distributions, i.e., without the need to estimate and evaluate a pdf. The other mentioned metrics require

integration over the domain of the pdfs, which will give computational issues since the considered pdfs will have a high dimensionality.

III. SCENARIO GENERATION

To generate realistic scenarios for the assessment of AVs, we use a data-driven approach: observed scenarios are used to generate new scenarios. To do this, the scenarios are parameterized, i.e., parameters are defined that characterize a scenario. For example, the duration of a scenario could be a parameter. Next, the pdf of the parameters is estimated. This pdf can be used to generate parameter values for new scenarios. In addition, the pdf contains the statistical information of the parameters so that the performance of AVs can be estimated [10], [32]. Choosing the parameters that describe a scenario, however, is not trivial:

- Choosing too few parameters might lead to an oversimplification of the actual scenarios. As a result, not all possible variations of a scenario are modeled.
- Too many parameters lead to problems with estimating the pdf, due to the curse of dimensionality [33].

To overcome this problem, we first consider as many parameters as needed for a complete description of the scenarios to avoid the oversimplification of the scenarios. Next, using an SVD, a new set of parameters is created using a linear mapping of the original scenario parameters. Because this new set of parameters is ordered according to the contribution of each of these parameters in describing the variation that exists among the original scenario parameters, we will consider only the most important parameters without losing too much information. In this way, the curse of dimensionality is avoided without relying on a predetermined choice of parameters.

Below, we first explain how to describe a scenario using many parameters. Next, Section III-B proposes the use of the SVD to reduce the number of parameters. Section III-C describes how KDE is used to estimate the pdf of the reduced set of parameters and how the estimated KDE can be used to generate new scenarios.

A. Parameterization of Scenarios

The first step of our approach is the parameterization of scenarios. There is no single best way to parameterize the scenarios considering the wide variety of scenarios. To deal with this variety, this work distinguishes quantitative scenarios from qualitative scenarios, using the definitions of *scenario* and *scenario category* of [34]:

Definition 1 (Scenario): A scenario is a quantitative description of the relevant characteristics and activities and/or goals of the ego vehicle(s), the static environment, the dynamic environment and all events that are relevant to the ego vehicle(s) within the time interval between the first and last relevant event.

Definition 2 (Scenario category): A scenario category is a qualitative description of relevant characteristics and activities and/or goals of the ego vehicle(s), the static environment, and the dynamic environment.

A scenario category is an abstraction of a scenario and, therefore, a scenario category comprises multiple scenarios [34]. For example, the scenario category “cut-in” comprises all possible cut-in scenarios. The goal of our approach is to determine the optimal parameterization of scenarios of a given scenario category based on a set of observed scenarios of the same scenario category and to estimate the pdf of these parameters that can be used to generate parameter values for new scenarios.

The observed scenarios are described using a time series for the content of the scenario that changes within the time window of the scenario (e.g., the speed of a vehicle) and some additional parameters for the content that is fixed (e.g., the lane width and the duration of the scenario). Here, $y(t) \in \mathbb{R}^{n_y}$ denotes the time series of a scenario with $t \in [t_0, t_1]$, where n_y denotes the dimension of the time series and t_0 and t_1 denote the start and end time of the scenario, respectively. The n_θ additional parameters are represented by $\theta \in \mathbb{R}^{n_\theta}$.

To deal with the time series, the continuous time interval $[t_0, t_1]$ is discretized, such that two consecutive time instants are $(t_1 - t_0)/(n_t - 1)$ apart. This gives:

$$\mathbf{y} = \begin{bmatrix} y(t_0) \\ y\left(t_0 + \frac{t_1 - t_0}{n_t - 1}\right) \\ y\left(t_0 + 2\frac{t_1 - t_0}{n_t - 1}\right) \\ \vdots \\ y(t_1) \end{bmatrix} \in \mathbb{R}^{n_t n_y}. \quad (1)$$

Note that n_t must be chosen such that no important information is lost during the discretization. Because in practice, due to the discrete nature of sensor readings, the time series $y(t)$ is obtained at certain specific times rather than on a continuous time interval, it may be required to use interpolation techniques, such as splines [35], to evaluate $\mathbf{y}$.

Let us assume that N_x observed scenarios can be used to generate new scenarios. To indicate that the scenario parameters $\mathbf{y}$ and θ belong to a specific scenario, the index $i \in \{1, \dots, N_x\}$ is used, i.e., the parameters of the i -th scenario are $\mathbf{y}_i$ and θ_i . To further ease the notation, $\mathbf{y}_i$ and θ_i are combined into one vector $\mathbf{x}_i$:

$$\mathbf{x}_i = \begin{bmatrix} \mathbf{y}_i \\ \theta_i \end{bmatrix} \in \mathbb{R}^{n_t n_y + n_\theta}. \quad (2)$$

B. Parameter Reduction Using Singular Value Decomposition

As shown in (2), $n_x = n_t n_y + n_\theta$ parameters describe a scenario. Even for small numbers of n_t , n_y , and n_θ , the total number of parameters becomes too large to reliably estimate the joint pdf. One way to avoid this curse of dimensionality is to assume that the parameters are independent, but especially the parameters $y(t_0)$ till $y(t_1)$ in (1) are obviously correlated, so assuming that the parameters are independent is not a good solution.

In the field of machine learning, Principal Component Analysis (PCA) is commonly used for dimensionality reduction [36]. As PCA uses the SVD [12], this work uses the SVD to transform the parameters $\mathbf{x}_i$ into a lower-dimensional

vector of parameters. Before applying the SVD, the parameters are weighted with $\alpha \in \mathbb{R}^{n_x}$ in order to give more or less importance to the n_x parameters. This is particularly useful to compensate for the imbalance in the parameter vector, where the imbalance is caused by the fact that the parameter vector considers the time series $y(t)$ at n_t different times and the additional parameters θ only once. Let us define a matrix that contains the parameters of the N_x scenarios:

$$X = [(\alpha \odot \mathbf{x}_1) - \mu \quad \cdots \quad (\alpha \odot \mathbf{x}_{N_x}) - \mu] \in \mathbb{R}^{n_x \times N_x}, \quad (3)$$

where $\odot$ denotes the element-wise product of vectors and $\mu \in \mathbb{R}^{n_x}$ denotes the mean of the weighted scenario parameters:

$$\mu = \frac{1}{N_x} \sum_{i=1}^{N_x} \alpha \odot \mathbf{x}_i. \quad (4)$$

Using the SVD of X , we obtain:

$$X = U \Sigma V^T. \quad (5)$$

Here, both $U \in \mathbb{R}^{n_x \times n_x}$ and $V \in \mathbb{R}^{N_x \times N_x}$ are orthonormal matrices. Therefore, both matrices can be interpreted as rotation matrices in $\mathbb{R}^{n_x}$ and $\mathbb{R}^{N_x}$, respectively. The matrix $\Sigma \in \mathbb{R}^{n_x \times N_x}$ takes the same shape as X . This matrix has only zeros except on (part of) the diagonal. The diagonal contains the so-called singular values, denoted by σ_j with $j \in \{1, \dots, \bar{N}\}$, $\bar{N} = \min(n_x, N_x)$. These singular values are in decreasing order, i.e.,

$$\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_{\bar{N}} \geq 0. \quad (6)$$

Because of the decreasing singular values, rotating the matrix X from the left with U^T transforms the data to a new coordinate system such that the first coordinate has the largest variance compared to the other coordinates. This variance equals σ_1^2 . Similarly, the second largest variance equals σ_2^2 and lies on the second coordinate, etc. Because of the decreasing variance, the scenario parameters can be approximated using only the first d coordinates of the new coordinate system, as these d coordinates describe the majority of the variations. So, the scenario parameters of the i -th scenario are approximated by setting $\sigma_j = 0$ for $j > d$:

$$\alpha \odot \mathbf{x}_i = \mu + \sum_{j=1}^{\bar{N}} \sigma_j v_{ij} \mathbf{u}_j \approx \mu + \sum_{j=1}^d \sigma_j v_{ij} \mathbf{u}_j, \quad (7)$$

where v_{ij} is the (i, j) -th element of V , $\mathbf{u}_j$ is the j -th column of U , and d is the number of parameters that are retained. Thus, the n_x parameters of the i -th scenario are approximated using the d parameters $v_{i1}, \dots, v_{id}$. The singular values $\sigma_1, \dots, \sigma_d$, the vectors $\mathbf{u}_1, \dots, \mathbf{u}_d$, and μ are used to map the new scenario parameters, $v_{i1}, \dots, v_{id}$, to an approximation of the weighted original scenario parameters, $\alpha \odot \mathbf{x}_i$.

Remark 1: Using the approximation of (7), it is not necessary to evaluate the complete SVD of (5). Only the first d columns of U and V need to be computed and only the first d singular values. In practice, $d \ll \bar{N}$, so this saves a significant amount of computation time.

The choice of $d < \bar{N}$ is not trivial. Choosing d too small results in too much loss of detail. Choosing d too large will give problems when estimating the pdf of the new parameters. One method to choose d is to look at the amount of overall variance of $\alpha \odot \mathbf{x}_i$ explained by the first d singular values. The overall variance scales with the sum of the squared singular values [12, p. 77], i.e.,

$$\sum_{i=1}^{N_x} ((\alpha \odot \mathbf{x}_i) - \mu)^T ((\alpha \odot \mathbf{x}_i) - \mu) = \sum_{j=1}^{\bar{N}} \sigma_j^2. \quad (8)$$

Thus, the first d singular values explain

$$\frac{\sum_{j=1}^d \sigma_j^2}{\sum_{j=1}^{\bar{N}} \sigma_j^2} \quad (9)$$

of the overall variance. One approach would be to set d such that (9) exceeds a certain threshold, such as 0.95. Another way to choose d is by inspecting the actual approximation error in (7) and keep increasing d until the approximation error is not too large. Section IV proposes an alternative way to determine d using a metric that quantifies the goal of our generated scenarios, i.e., that the generated scenarios are representing real-world scenarios and cover the actual variety of real-world scenarios.

C. Estimating the probability density function

Using the approximation of (7) based on the SVD, the i -th scenario is described by the vector $\tilde{\mathbf{v}}_i$:

$$\tilde{\mathbf{v}}_i^T = [v_{i1} \quad \cdots \quad v_{id}]. \quad (10)$$

Note that the d entries of $\tilde{\mathbf{v}}_i$ are linearly uncorrelated with the d entries of $\tilde{\mathbf{v}}_m$ ($m \neq i$).¹ Despite the linear independence, the different entries of $\tilde{\mathbf{v}}_i$ may still be dependent due to higher-order correlations; so we treat these d entries as dependent variables.

To estimate the pdf of $\tilde{\mathbf{v}}_i$, we propose to use KDE. KDE [13], [14] is often referred to as a non-parametric way to estimate the pdf, because KDE does not rely on the assumption that the data are drawn from a given parametric family of probability distributions. Because KDE produces a pdf that adapts itself to the data, it is flexible regarding the shape of the actual underlying distribution of $\tilde{\mathbf{v}}_i$. In KDE, the pdf is estimated as:

$$\hat{f}_H(v) = \frac{1}{N_x} \sum_{i=1}^{N_x} K_H(v - \tilde{\mathbf{v}}_i). \quad (11)$$

Here, $K_H(\cdot)$ is the so-called scaled kernel with a positive definite symmetric bandwidth matrix $H \in \mathbb{R}^{d \times d}$. The kernel $K(\cdot)$ and the scaled kernel $K_H(\cdot)$ are related using

$$K_H(u) = |H|^{-1/2} K(H^{-1/2}u), \quad (12)$$

¹This is assuming that $\sigma_d > 0$. With this assumption and because X in (3) is defined such that the sum of each row of X equals zero, it is easy to verify that $\frac{1}{N_x} \sum_{m=1}^{N_x} v_{mj} = 0$ for $j \in \{1, \dots, N_x\}$. Therefore, $\sum_{i=1}^{N_x} (v_{ij} - \frac{1}{N_x} \sum_{m=1}^{N_x} v_{mj}) (v_{ik} - \frac{1}{N_x} \sum_{m=1}^{N_x} v_{mk}) = \sum_{i=1}^{N_x} v_{ij} v_{ik} = 0$ for $j \neq k$, where the latter equality follows from the orthonormality of V .

where $|\cdot|$ denotes the matrix determinant. The choice of the kernel function is not as important as the choice of the bandwidth matrix [37], [38]. This article considers the Gaussian kernel,² which is given by

$$K(u) = \frac{1}{(2\pi)^{d/2}} \exp\left\{-\frac{1}{2}\|u\|_2^2\right\}, \quad (13)$$

where $\|u\|_2^2 = u^T u$ denotes the squared 2-norm of u .

A bandwidth matrix of the form $H = h^2 I_d$ is used, where I_d denotes the d -by- d identity matrix. The bandwidth h is determined with leave-one-out cross-validation [41], because this minimizes the difference between the real pdf and the estimated pdf according to the Kullback-Leibler divergence [37], [42].

To sample scenario parameters using $\hat{f}_H(\cdot)$, first, an integer $i \in \{1, \dots, N_x\}$ is randomly chosen with each integer having equal likelihood. Next, a random sample is drawn from a Gaussian with covariance H and mean $\tilde{\mathbf{v}}_i$. Then, using the approximation in (7), the scenario parameters are calculated.

As far as the computational effort is concerned, sampling the scenario parameters from a KDE is efficient because there is no need to actually evaluate the pdf. Determining the optimal bandwidth matrix requires more computational effort, but this only has to be done once per data set. The computational complexity of cross-validation methods for the bandwidth estimation typically scales with N_x^2 [43].

IV. SCENARIO REPRESENTATIVENESS METRIC

Ideally, the parameters of the generated scenarios are sampled from the same distribution that underlies the real-world scenario parameters. The problem is that this distribution is unknown. Nevertheless, it is possible to define a metric that quantifies the similarity of the distribution that is used to generate scenario parameters and the distribution that underlies the real-world scenario parameters. Section IV-A further explains the goal of this metric, which we call the SR metric. Next, Section IV-B explains the Wasserstein distance [15], which is then applied to derive our metric in Section IV-C.

A. Scenario Comparison Problem

The set of observed scenarios, described using the parameters $\mathbf{x}_i$, $i \in \{1, \dots, N_x\}$, are used for generating the scenario parameters. To ease the notation, let us denote the set of observed scenarios by $\mathcal{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_{N_x}\}$. This work assumes that these scenarios — that are comprised by the same scenario category — are independently and identically distributed according to the distribution $f(\cdot) : \mathbb{R}^{n_x} \rightarrow \mathbb{R}$. Let us denote the set of generated scenario parameter vectors by $\mathcal{W} = \{\mathbf{w}_1, \dots, \mathbf{w}_{N_w}\}$ where $\mathbf{w}_i \in \mathbb{R}^{n_x}$, $i \in \{1, \dots, N_w\}$ are similarly parameterized as in (2) and N_w is the number of generated scenario parameter vectors. Let $\hat{f}(\cdot) : \mathbb{R}^{n_x} \rightarrow \mathbb{R}$ denote the pdf of the generated scenario parameter vectors,

²The advantage of the Gaussian kernel is that it gives the possibility to calculate a metric that quantifies the completeness of the data [39] and to apply conditional sampling when generating scenario parameters [40]. Both these topics are out of scope of this article.

which is obtained from $\hat{f}_H(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}$ under a change of variable according to the approximation in (7). As later appears, it is not needed to have an explicit definition for $\hat{f}(\cdot)$. Ideally, $\hat{f}(\cdot)$ is equal to $f(\cdot)$. So our metric aims to quantify the similarity of $\hat{f}(\cdot)$ and $f(\cdot)$.

To estimate the similarity between $\hat{f}(\cdot)$ and $f(\cdot)$, we cannot simply compare $\mathcal{W}$ with $\mathcal{X}$. In that case, taking $\mathcal{W} = \mathcal{X}$ would give us the best result, but this is undesirable because, ideally, the scenarios of the generated parameters cover the whole variety of real-world scenarios and not just the variety that have been observed in $\mathcal{X}$. Therefore, another set of scenarios is needed that can be used to test. Let us assume that such a set of scenarios is available, denoted by $\mathcal{Z} = \{\mathbf{z}_1, \dots, \mathbf{z}_{N_z}\}$ where $\mathbf{z}_i \in \mathbb{R}^{n_x}$, $i \in \{1, \dots, N_z\}$ are independently and identically distributed according to $f(\cdot)$. Thus, $\mathcal{X}$ and $\mathcal{Z}$ can be regarded as a training and test set, respectively.

In summary, the goal is to find a metric that quantifies the similarity of $\hat{f}(\cdot)$ and $f(\cdot)$ using the sets of observed scenario parameters $\mathcal{X}$ and $\mathcal{Z}$ and the set of scenario parameters $\mathcal{W}$, generated based on $\mathcal{X}$.

B. Empirical Wasserstein Metric

The p -th Wasserstein metric ($p \geq 1$) [15] is used to compare two pdfs $\zeta(\cdot)$ and $\eta(\cdot)$ defined on the set $\mathcal{U}$. This metric is defined as follows:

$$W_p(\zeta, \eta) = \left(\inf_{\gamma \in \Gamma(\zeta, \eta)} \left\{ \int_{\mathcal{U} \times \mathcal{U}} (\Delta(u, v))^p d\gamma(u, v) \right\} \right)^{1/p}. \quad (14)$$

Here, $\Delta(u, v)$ denotes the distance from u to v , which will be defined below, and $\Gamma(\zeta, \eta)$ denotes the set of joint distributions of (u, v) that have marginal distributions $\zeta(\cdot)$ and $\eta(\cdot)$. Intuitively, if the pdfs $\zeta(\cdot)$ and $\eta(\cdot)$ are seen as two piles of earth having a different shape with mass 1, then (14) calculates the minimum cost of converting one pile of earth with shape $\zeta(\cdot)$ into a pile of earth with shape $\eta(\cdot)$. Therefore, the Wasserstein metric is also referred to as the earth mover's distance [44].

In our case, the goal is to have a metric to compare $f(\cdot)$ and $\hat{f}(\cdot)$. Because $f(\cdot)$ is unknown, its approximation based on $\mathcal{Z}$ is considered:

$$f(\mathbf{z}) \approx \frac{1}{N_z} \sum_{i=1}^{N_z} \delta(\mathbf{z} - \mathbf{z}_i), \mathbf{z} \in \mathbb{R}^{n_x}, \quad (15)$$

where $\delta(\cdot)$ denotes the Dirac delta function. Considering the high dimension of $\mathbf{z}$, numerical approximation of the integral of the Wasserstein metric (14) using this approximation and $\hat{f}(\cdot)$ would require so many evaluations of $\hat{f}(\cdot)$ that it becomes computationally infeasible. Therefore, the empirical estimation of the Wasserstein metric (14) is considered, which makes use of the empirical estimation of $\hat{f}(\cdot)$:

$$\hat{f}(\mathbf{w}) \approx \frac{1}{N_w} \sum_{i=1}^{N_w} \delta(\mathbf{w} - \mathbf{w}_i), \mathbf{w} \in \mathbb{R}^{n_x}. \quad (16)$$

Substituting the empirical estimations of (15) and (16) for $\zeta(\cdot)$ and $\eta(\cdot)$, respectively, into (14), leads to the so-called

empirical Wasserstein metric [45], which is defined as:

$$\tilde{W}_p(\mathcal{Z}, \mathcal{W}) = \left(\inf_T \sum_{i=1}^{N_z} \sum_{j=1}^{N_w} (\Delta(\mathbf{z}_i, \mathbf{w}_j))^p T_{ij} \right)^{1/p}, \quad (17)$$

where T_{ij} is the (i, j) -th element of the transportation matrix T that is subject to the following conditions:

$$\sum_{i=1}^{N_z} T_{ij} = \frac{1}{N_w} \quad \forall j \in \{1, \dots, N_w\}, \quad (18)$$

$$\sum_{j=1}^{N_w} T_{ij} = \frac{1}{N_z} \quad \forall i \in \{1, \dots, N_z\}, \quad (19)$$

$$T_{ij} \geq 0 \quad \forall i \in \{1, \dots, N_z\}, j \in \{1, \dots, N_w\} \quad (20)$$

For the distance function, we will use the 2-norm of the difference of the scenario parameters after scaling the scenario parameters according to the weights α that we also used in Section III-B:

$$\Delta(\mathbf{z}, \mathbf{w}) = \|(\alpha \odot \mathbf{z}) - (\alpha \odot \mathbf{w})\|_2. \quad (21)$$

C. Metric for Testing Scenario Representativeness

The empirical Wasserstein metric $\tilde{W}_p(\mathcal{Z}, \mathcal{W})$ is an approximation of the Wasserstein metric $W_p(f, \hat{f})$. As one might expect, using an infinite number of scenario parameters, i.e., for $N_z \rightarrow \infty$ and $N_w \rightarrow \infty$, the empirical Wasserstein metric approaches the Wasserstein metric with probability 1 [45]. The problem is that N_z and N_w are not infinite. In addition, whereas a fairly large number for N_w can be chosen, as it is only limited by the available computational resources, to increase N_z , more data are needed and this is generally expensive. Therefore, this work proposes a metric that is different from (17).

Our proposed SR metric is based on the following intuition: Suppose that $\hat{f}$ is indeed an approximation of f . Because $\mathcal{X}$ and $\mathcal{Z}$ are based on the same underlying pdf, i.e., f , it is expected that $\tilde{W}_p(\mathcal{X}, \mathcal{W})$ is similar to $\tilde{W}_p(\mathcal{Z}, \mathcal{W})$. If, however, $\tilde{W}_p(\mathcal{X}, \mathcal{W})$ is significantly smaller than $\tilde{W}_p(\mathcal{Z}, \mathcal{W})$, it suggests overfitting of the training data because the generated scenario parameters are too much skewed towards the training data $\mathcal{X}$. To penalize overfitting of the training data, our SR metric includes a penalty in case $\tilde{W}_p(\mathcal{Z}, \mathcal{W})$ is larger than $\tilde{W}_p(\mathcal{X}, \mathcal{W})$. Thus, the SR metric becomes:

$$M_p(\mathcal{W}, \mathcal{Z}, \mathcal{X}) = \tilde{W}_p(\mathcal{Z}, \mathcal{W}) + \beta \left(\tilde{W}_p(\mathcal{Z}, \mathcal{W}) - \tilde{W}_p(\mathcal{X}, \mathcal{W}) \right). \quad (22)$$

Here, β is the weight of the penalty. The case study in Section V demonstrates empirically that $M_p(\mathcal{W}, \mathcal{Z}, \mathcal{X})$ of (22) better correlates with the Wasserstein metric of (14) than the empirical Wasserstein metric of (17) and a method to choose β .

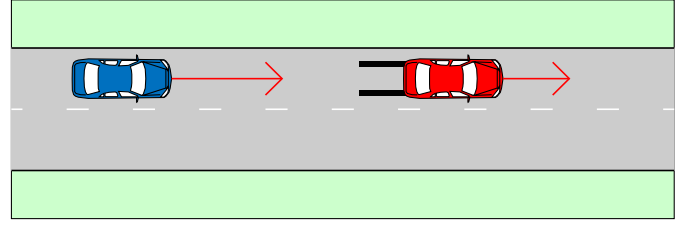


Fig. 1. Schematic representation of the scenario category “leading vehicle decelerating (LVD).” The left vehicle is the ego vehicle.

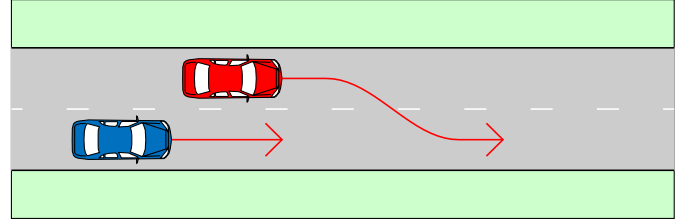


Fig. 2. Schematic representation of the scenario category “cut-in.” The left vehicle is the ego vehicle.

V. CASE STUDY

To illustrate the proposed method for generating the scenario parameters (Section III) and the SR metric (Section IV), these are applied in a case study. Section V-A explains the scenario categories that are considered in the case study and describes the choices that are made regarding the scenario parameterization. Section V-B illustrates the approximation of the original parameterization using the SVD. Next, the scenario parameter generation method is demonstrated in Section V-C. Section V-C also shows that the SR metric (22) can be used to choose d . Our method for generating scenario parameters is compared with other methods in Section V-D. Section V-E demonstrates that the SR metric (22) better correlates with the Wasserstein metric (14) than the empirical Wasserstein metric (17).

A. Scenario Categories and Parameterization

In this case study, two scenario categories are considered. The first scenario category, labeled leading vehicle decelerating (LVD), involves an ego vehicle that is following another vehicle that decelerates, see Fig. 1. As a result, the ego vehicle might need to brake or change direction to avoid contact with the vehicle that decelerates. The second scenario category considers a vehicle that performs a cut-in, such that this vehicle becomes the leading vehicle of the ego vehicle, see Fig. 2. Depending on the speed and timing of the vehicle that performs a cut-in, the ego vehicle might need to brake or change direction to avoid a collision.

To obtain the scenarios, the data set described in [46] is used. The data were recorded from a single vehicle in which 20 drivers were asked to drive a prescribed route, resulting in 63 hours of data containing 1150 LVD scenarios and 289 cut-in scenarios. The majority of the route was on the highway. To measure the surrounding traffic, the vehicle was equipped with three radars and one camera. The surrounding traffic was measured by fusing the data of the

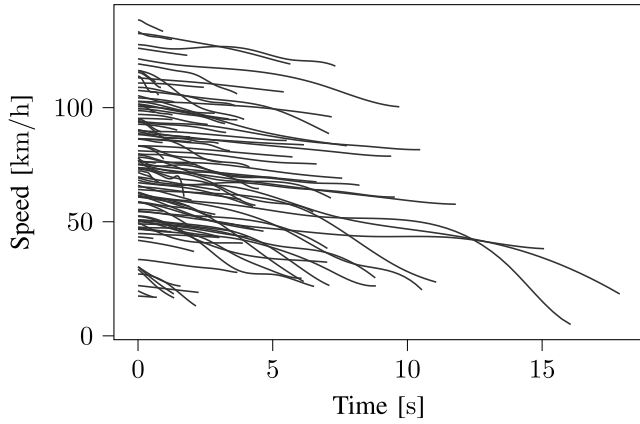


Fig. 3. Speed of the leading vehicle during 100 randomly-selected observed LVD scenarios. For plotting purposes, the starting time of each scenario is set to 0.

radars and the camera as described in [47]. To extract the LVD and cut-in scenarios from the data set with the fused data, we searched for particular (combinations of) activities in the data: a deceleration activity of a leading vehicle indicates an LVD scenario and a lane change of another vehicle that becomes the leading vehicle indicates a cut-in scenario. For more information on the process of extracting the scenarios, see [48].

From the 1150 LVD scenarios, the training uses 80% (so $N_x = 920$) and the testing uses the remaining 20% (so $N_z = 230$) as this 80/20 ratio is commonly used for splitting the data into a training set and a test set. The training data are used for generating $N_w = 10000$ new scenario parameter vectors. To describe the decelerating behavior of the leading vehicle, the acceleration of the leading vehicle at $n_t = 50$ time instants is used ($n_y = 1$). As additional parameters, the duration of the scenario, $t_1 - t_0$, the initial speed of the leading vehicle, and the initial time gap between the leading vehicle and the ego vehicle are considered ($n_\theta = 3$). Thus, $n_x = 53$. In Fig. 3, the speed of the leading vehicle of 100 randomly-selected observed LVD scenarios are shown. The k -th weight, α_k , is obtained by dividing a chosen constant β_k by the standard deviation of the k -th parameter:

$$\alpha_k = \frac{\beta_k}{\sqrt{\frac{1}{N_x} \sum_{i=1}^{N_x} ((\mathbf{x}_i)_k - \bar{\mathbf{x}}_k)^2}}, \quad (23)$$

with $(\mathbf{x}_i)_k$ denoting the k -th element of $\mathbf{x}_i$ and $\bar{\mathbf{x}}_k = \frac{1}{N_x} \sum_{i=1}^{N_x} (\mathbf{x}_i)_k$. In this way, the contribution of the k -th parameter to the overall variance (see (8)) only depends on β_k . When choosing $\beta_1 = \dots = \beta_{53}$, the acceleration of the leading vehicle would contribute 50 times more to the overall variance of (8), because $n_t = 50$ elements are used to describe the acceleration. For the LVD scenarios, we want to give the acceleration the same importance as each of the other parameters, so we choose $\beta_1 = \dots = \beta_{50} = 1/\sqrt{n_t}$ and $\beta_{51} = \beta_{52} = \beta_{53} = 1$.

From the 289 cut-in scenarios, 80% are used for training (so $N_x = 231$) and 20% are used for testing (so $N_z = 58$). Both cut-in scenarios from the left and from the right are considered. The training data are used for generating $N_w =$

TABLE I
THE LAST $n_\theta = 3$ COORDINATES OF BOTH μ AND THE FIRST FOUR COLUMNS OF U AFTER SCALING WITH α FOR THE LVD SCENARIOS. NOTE THAT $\odot$ DENOTES ELEMENT-WISE DIVISION

	Coordinate 51 Scenario duration	Coordinate 52 Initial speed	Coordinate 53 Initial time gap
$\mu \odot \alpha$	4.73 s	22.11 km/h	1.49 s
$\mathbf{u}_1 \odot \alpha$	-1.50 s	-15.17 km/h	0.28 s
$\mathbf{u}_2 \odot \alpha$	-3.09 s	12.22 km/h	-0.06 s
$\mathbf{u}_3 \odot \alpha$	1.15 s	-16.52 km/h	0.29 s
$\mathbf{u}_4 \odot \alpha$	1.32 s	-2.88 km/h	-0.08 s

10000 parameter vectors that describe cut-in scenarios. A cut-in scenario is described using the speed of the vehicle that performs the lane change and its lateral position with respect to the center of the ego vehicle's lane (so $n_y = 2$) at $n_t = 50$ time instants. In case of a cut-in scenario from the left, the lateral position is positive when the cutting-in vehicle is on the left of the center of ego vehicle's lane and vice versa for a cut-in scenario from the right. Furthermore, $n_\theta = 3$ extra parameters are used to describe a cut-in scenario: the duration of the scenario, the initial speed of the ego vehicle, and the initial longitudinal position of the cutting-in vehicle with respect to the ego vehicle. Thus, $n_x = 103$. To give the same importance to the speed of the vehicle that performs the lane change, its lateral position, and the 3 extra parameters, the weights are calculated using (23) with $\beta_1 = \dots = \beta_{100} = 1/\sqrt{n_t}$ and $\beta_{101} = \beta_{102} = \beta_{103} = 1$.

B. Approximation of Scenarios With SVD

As explained in Section III-C, using too many parameters will lead to poor estimations of the pdf of the parameters. We use an SVD to obtain a reduced number of parameters that best describe the original scenarios parameters. This section illustrates the approximation of the original scenario parameters using the parameters obtained after applying the SVD.

Following the approximation of (7), the scaled parameter vector, $\alpha \odot \mathbf{x}_i$, is approximated using a linear combination of the first d columns of U , i.e., $\mathbf{u}_1, \dots, \mathbf{u}_d$. In Fig. 4 and Table I, μ and the first four columns of U are shown for the LVD scenarios. For an easier interpretation, the original scaling of the parameters by α is undone via the element-wise division by α . Fig. 4 shows that the average scenario starts with a deceleration of about 0.4 m/s² and ends with a deceleration of about 0.8 m/s². Table I shows that the average scenario duration is 4.73 s, the average initial speed of the leading vehicle is 22.11 km/h, and the average initial time gap is 1.49 s. Since each scenario is estimated by combining the curves in Fig. 4 and values in Table I, it can be seen that the approximations do not contain complex acceleration curves. In other words, the accelerations will be smoothed and the details may get lost. The amount of smoothing depends on d , i.e., the number of vectors of U that are used to approximate the original parameter vector. Choosing the value of d is a trade-off: a higher value of d leads to less smoothing and, therefore, a smaller approximation error, but choosing d too

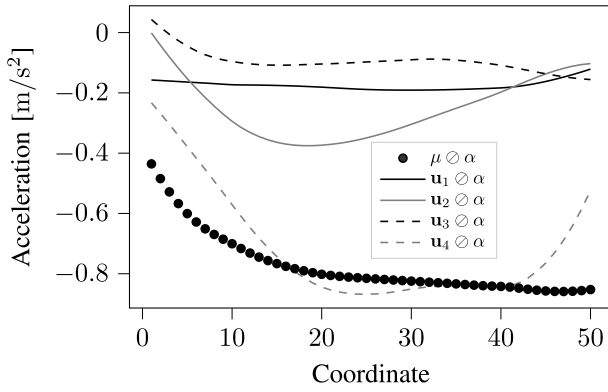


Fig. 4. The first $n_t = 50$ coordinates of both μ and the first four columns of U after scaling with α for the LVD scenarios. Note that $\otimes$ denotes element-wise division.

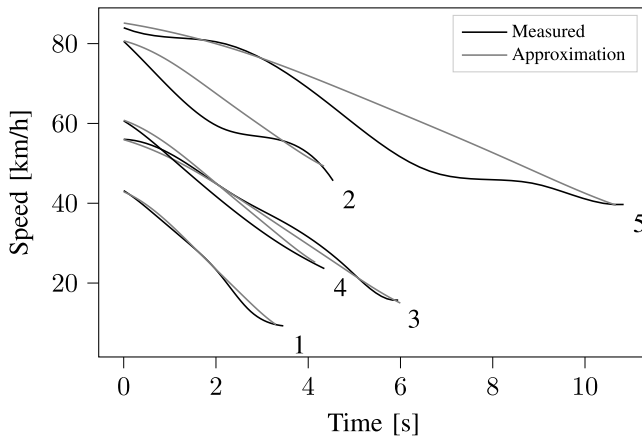


Fig. 5. Five scenarios that require the highest average deceleration of the follower. The black lines denote the observed scenarios and the gray lines denote their approximations based on the $d = 4$ parameters. The corresponding initial time gaps are listed in Table II.

large leads to problems when estimating the pdf of the new parameters.

Fig. 5 shows five LVD scenarios. These selected LVD scenarios correspond to the five LVD scenarios that require the highest average deceleration of the following vehicle. The line with the “1” denotes the LVD scenario that requires the highest average deceleration. Table II lists the values of $\sigma_j v_{ji}$ for $j \in \{1, \dots, d\}$ with $d = 4$ that are used to approximate the original scenarios according to the approximation in (7). The gray lines in Fig. 5 show the approximated speed of the five LVD scenarios. Table II shows the initial time gaps of the five scenarios shown in Fig. 5. These five scenarios illustrate that the accelerations are smoothed, but the main characteristics of the scenarios are captured by the approximations: the average deceleration, the scenario duration, the initial speed, and the initial time gap are well approximated.

C. Generating Scenario Parameters

An important parameter for the generation of the scenario parameter vectors is the number of reduced parameters (d). One approach is to look at the so-called explained variance of (9) of the first d singular values, see Table III. The first 4 singular values already explain 90.4% of the variance for

TABLE II
INITIAL TIME GAPS OF THE FIVE SCENARIOS THAT REQUIRE THE HIGHEST AVERAGE DECELERATION OF THE FOLLOWER. THE CORRESPONDING SPEEDS ARE SHOWN IN FIG. 5

#	$\sigma_1 v_{i1}$	$\sigma_2 v_{i2}$	$\sigma_3 v_{i3}$	$\sigma_4 v_{i4}$	Initial time gap	
					Original	Approximated
1	2.21	0.30	-0.03	2.16	1.91 s	1.91 s
2	-0.28	0.75	-0.46	1.56	1.08 s	1.11 s
3	0.61	-0.21	-0.42	1.53	1.43 s	1.43 s
4	1.74	0.25	0.58	1.63	2.00 s	2.00 s
5	-1.74	-0.71	-0.49	1.30	0.81 s	0.80 s

TABLE III
EXPLAINED VARIANCE ACCORDING TO (9)

d	Leading vehicle decelerating	Cut-in
1	36.9%	36.9%
2	63.0%	63.3%
3	78.0%	84.5%
4	90.4%	94.1%
5	94.5%	96.9%
6	96.7%	99.0%
7	98.2%	99.6%
8	99.2%	99.8%

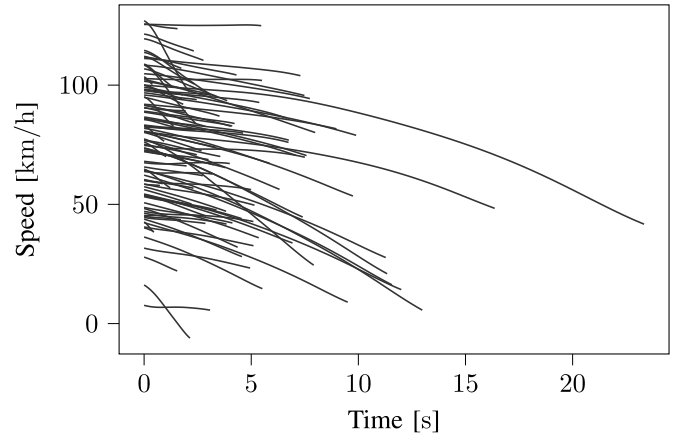


Fig. 6. Speed of the leading vehicle during 100 generated LVD scenarios.

the LVD scenarios, so $d = 4$ might be a suitable choice. In Fig. 6, the speed of the leading vehicle of 100 generated LVD scenarios is shown using $d = 4$.

Another way to determine d is to use the SR metric $M_p(\mathcal{W}, \mathcal{Z}, \mathcal{X})$ defined in (22). In Fig. 7, the result is shown when applying this metric with $p = 1$, alongside with the empirical Wasserstein metric $\tilde{W}_1(\mathcal{Z}, \mathcal{W})$ and the penalty $\tilde{W}_1(\mathcal{Z}, \mathcal{W}) - \tilde{W}_1(\mathcal{X}, \mathcal{W})$. Each point in Fig. 7 represents the median³ when applying the metric 200 times, each time with a different (random) partition of the training data $\mathcal{X}$ and test data $\mathcal{Z}$. The standard deviation of the medians in Fig. 7, estimated using bootstrapping [50], is 0.005 or less. For the SR metric, the penalty is weighted using $\beta = 0.25$. The choice of $\beta = 0.25$ is justified in Section V-E.

The most left points in Fig. 7 represent the metric in case the set of training data $\mathcal{X}$ is directly used to sample the scenario parameters instead of the approach of Section III. Here, $\mathcal{W}$ is

³We preferred to use the median instead of the mean, such that the result is less influenced by outliers [49].

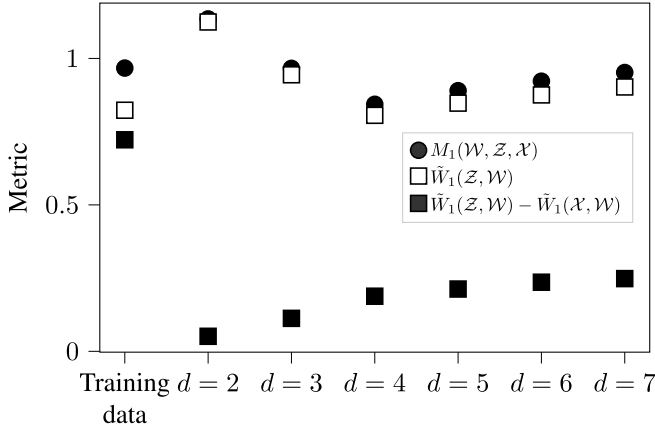


Fig. 7. Medians of the metrics for the set of generated LVD scenario parameters. Note that $d = 1$ is excluded, because its metrics are an order of magnitude higher than for $d = 2$ and would, therefore, not be visible with the current scaling of the y-axis.

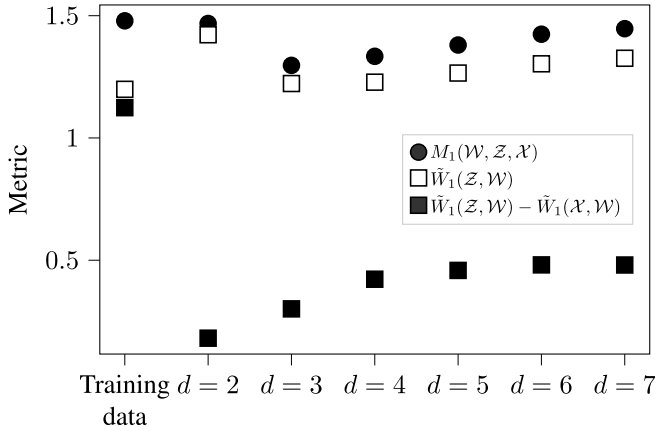


Fig. 8. Medians of the metrics for the set of generated cut-in scenario parameters.

a selection with replacement of N_w scenarios from $\mathcal{X}$, i.e.:

$$\mathbf{w}_i = \mathbf{x}_{\lfloor u \rfloor}, u \sim U(1, N_x + 1), \quad \forall i \in \{1, \dots, N_w\}, \quad (24)$$

where $U(1, N_x + 1)$ denotes the continuous uniform distribution with boundaries 1 and $N_x + 1$, and $\lfloor \cdot \rfloor$ denotes the floor function. Using the training data directly for “generating scenarios” leads to a low empirical Wasserstein metric. The downside is that there is not much variation among the generated scenarios. Therefore, the penalty is also the highest, which results in $M_1(\mathcal{W}, \mathcal{Z}, \mathcal{X}) \approx 0.967$. Looking at $d = 4$, the empirical Wasserstein metric (open squares) is approximately similar compared to when the training set is directly used. Due to the sampling of the scenario parameters from the KDE, the generated scenarios contain more variation than the training set, resulting in a lower penalty and, therefore, a lower metric evaluation of $M_1(\mathcal{W}, \mathcal{Z}, \mathcal{X}) \approx 0.843$. Increasing d even further results in higher metric evaluations. So based on the proposed metric, $d = 4$ seems the right choice.

Fig. 8 shows the results of the generation of the cut-in scenario parameters in a similar way as Fig. 7. The standard deviation of all points in Fig. 8 is less than 0.008. The lowest penalty is obtained with $d = 2$, but the higher empirical Wasserstein distance suggests that too much information is

TABLE IV
MEDIAN OF THE METRIC $M_1(\mathcal{W}, \mathcal{Z}, \mathcal{X})$ WITH DIFFERENT APPROACHES FOR GENERATING SCENARIOS

Parameters	Distribution	Dependency	LVD	Cut-in
SVD	KDE	Dependent	0.84	1.30
SVD	Gaussian	Dependent	1.00	1.33
SVD	KDE	Independent	0.99	1.28
SVD	Gaussian	Independent	1.00	1.33
Fixed	KDE	Dependent	2.65	1.70
Fixed	Gaussian	Dependent	2.58	1.71
Fixed	KDE	Independent	2.31	1.67
Fixed	Gaussian	Independent	4.76	1.69

lost. The best result, i.e., where the SR metric, $M_1(\mathcal{W}, \mathcal{Z}, \mathcal{X})$, is minimal, is obtained at $d = 3$.

D. Comparison With Other Approaches

Our proposed method utilizes an SVD to obtain the scenario parameters and multivariate KDE to estimate the pdf of these parameters. To illustrate the advantages of these choices, the results of our method are compared with alternative approaches. First, instead of using an SVD for obtaining the parameters, a fixed parameterization is used, such as in [10], [18], [21]. Second, instead of using KDE to estimate the pdf of the parameters, a Gaussian distribution like in [29] is assumed. Third, the parameters are assumed to be independent.

When using a fixed parameterization for the LVD scenario, 4 parameters describe the scenario [10]: the speed reduction of the leading vehicle, the final speed of the leading vehicle, the duration of the scenario, and the initial time gap between the leading vehicle and the ego vehicle. The speed of the leading vehicle is assumed to follow a sinusoidal function, such that the acceleration at the start and at the end of the scenario equals zero. In case of the cut-in scenario, 5 parameters describe the scenario: the mean speed of the vehicle cutting in, its initial lateral position with respect to the center of the ego vehicle’s lane, the duration of the scenario, the initial speed of the ego vehicle, and the initial longitudinal position of the vehicle cutting in with respect to the ego vehicle. The speed of the vehicle cutting in is assumed to be constant. Its lateral position is assumed to follow a sinusoidal function, such that the vehicle ends at the center of the ego vehicle’s lane. For estimating the pdf of these parameters, the comparison considers 4 possibilities: multivariate KDE, multiple univariate KDEs, a multivariate Gaussian distribution, and multiple univariate Gaussian distributions.

Table IV shows the results of the different approaches for generating scenario parameters. For the LVD scenarios, our proposed approach (top row in Table IV) resulted in the lowest $M_1(\mathcal{W}, \mathcal{Z}, \mathcal{X})$. For the cut-in scenarios, it is interesting to note that the scores are not very different if SVD is used to obtain the parameters. This is partly explained by the smaller data set, because this results in a higher bandwidth⁴ that makes the KDE result with the Gaussian kernel look more like a Gaussian distribution. Using SVD and KDE while assuming that the

⁴On average, the bandwidth is about 1.5 to 2 times larger for the cut-in scenarios compared to the LVD scenarios.

parameters are independent, results in an even better result: 1.28 instead of 1.30 (with a standard deviation of 0.005). This indicates that assuming that the 3 parameters obtained with the SVD are independent, is acceptable.

E. Evaluating the Scenario Representativeness Metric

To determine whether our proposed metric (22) correlates better with the Wasserstein metric (14) than the empirical Wasserstein metric (15), the Wasserstein metric (14) needs to be known. This is not possible because the true underlying distribution of the data is unknown. To estimate the Wasserstein metric (14), the empirical Wasserstein metric (15) can be used with large numbers of test scenarios and generated scenario parameters, i.e., with large values of N_z and N_w , respectively. Since a large number of test scenarios is not available to us, we assume a certain distribution for $f(\cdot)$ from which the training data and the test data are generated. The approach is as follows (the numbers are for the LVD scenarios and, in parenthesis, the numbers for the cut-in scenarios are shown):

- 1) Based on the original 1150 (289) scenarios, obtained from the data, the following sets of scenario parameters are generated using the proposed approach explained in Section III with $d = 4$ ($d = 3$):
 - A new set of training data $\mathcal{X}^*$ of size $N_x = 920$ ($N_x = 231$);
 - A new set of test data $\mathcal{Z}^*$ of size $N_z = 230$ ($N_z = 58$); and
 - A large set of test data $\mathcal{Z}_{\text{large}}^*$ of size $N_z = 10000$ ($N_z = 10000$).
- 2) Based on $\mathcal{X}^*$, $N_w = 10000$ ($N_w = 10000$) scenario parameters are generated and collected in a set $\mathcal{W}^*$.
- 3) Our proposed metric is computed using $\mathcal{W}^*$, $\mathcal{Z}^*$, and $\mathcal{X}^*$: $M_1(\mathcal{W}^*, \mathcal{Z}^*, \mathcal{X}^*)$ with $\beta = 0.25$.
- 4) The Wasserstein metric of (14) is estimated using the empirical Wasserstein metric of (17) with $\mathcal{W}^*$ and $\mathcal{Z}_{\text{large}}^*$. Note: to approximate the Wasserstein metric of (14) using the empirical Wasserstein metric of (17), both $\mathcal{W}^*$ and $\mathcal{Z}_{\text{large}}^*$ need to be large (but not necessarily the same) in size.

We have repeated this approach 200 times, each time with a different (random) partition of the training data $\mathcal{X}$ and test data $\mathcal{Z}$. Figs. 9 and 10 show the result of this approach for the LVD scenarios and cut-in scenarios, respectively. In both cases, the empirical Wasserstein metric $\tilde{W}_1(\mathcal{Z}^*, \mathcal{W}^*)$ is minimal when the training data are directly used for the generated scenario parameters. Thus, the empirical Wasserstein metric suggests that the best approach for generating new scenario parameters is to simply sample parameters from the training data. The actual Wasserstein metric, estimated using $\tilde{W}_1(\mathcal{Z}_{\text{large}}^*, \mathcal{W}^*)$ shows that using our proposed method outperforms sampling parameters directly from the training data.

To justify the choice of $\beta = 0.25$, Fig. 11 shows the correlation between the medians of the proposed metric $M_1(\mathcal{W}^*, \mathcal{Z}^*, \mathcal{X}^*)$ and $\tilde{W}_1(\mathcal{Z}_{\text{large}}^*, \mathcal{W}^*)$ for different values of β . With $\beta = 0$, i.e., $M_1(\mathcal{W}^*, \mathcal{Z}^*, \mathcal{X}^*) = \tilde{W}_1(\mathcal{Z}^*, \mathcal{W}^*)$, the correlation is 0.974 for the LVD scenarios and 0.824 for

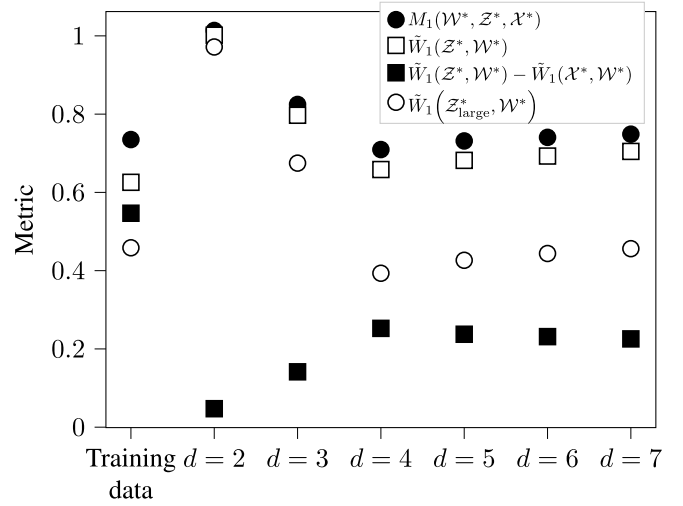


Fig. 9. Medians of the metrics for the set of $N_w = 10000$ generated LVD scenario parameter vectors. In this case, the $N_x = 920$ scenarios of $\mathcal{X}^*$ are sampled from $\hat{f}_H(\cdot)$ of (11), where $\hat{f}_H(\cdot)$ is based on the original data set $\mathcal{X}$.

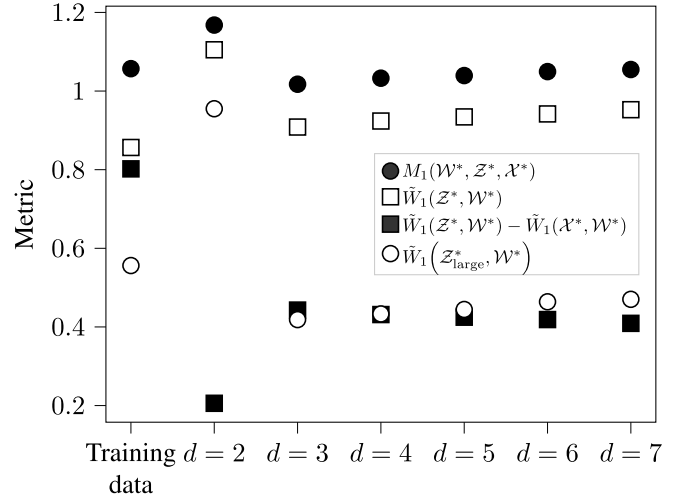


Fig. 10. Medians of the metrics for the set of $N_w = 10000$ generated cut-in scenario parameter vectors.

the cut-in scenarios. The correlation increases with increasing β until the maximum is obtained at $\beta \approx 0.21$ for the LVD scenarios and at $\beta \approx 0.27$ for the cut-in scenarios. The correlations at these maxima are 0.992 and 0.987, respectively. Increasing β further results in a lower correlation, which suggests that a choice of $\beta = 0.25$ seems appropriate.

The experiment described in this section can be used to determine both d and β in an iterative manner given an initial choice for β (denoted by β_0):

- 1) Set $i = 0$.
- 2) Determine d_i , i.e., the optimal number of parameters that minimizes $M_1(\mathcal{W}, \mathcal{Z}, \mathcal{X})$ using $\beta = \beta_i$.
- 3) Generate $\mathcal{X}^*$, $\mathcal{W}^*$, $\mathcal{Z}^*$, and $\mathcal{Z}_{\text{large}}^*$ using the approach described in this section with $d = d_i$.
- 4) Increase i by 1.
- 5) Determine β_i by maximizing the correlation between $M_1(\mathcal{W}^*, \mathcal{Z}^*, \mathcal{X}^*)$ with $\beta = \beta_i$ and $\tilde{W}_1(\mathcal{Z}_{\text{large}}^*, \mathcal{W}^*)$ (e.g., see Fig. 11).

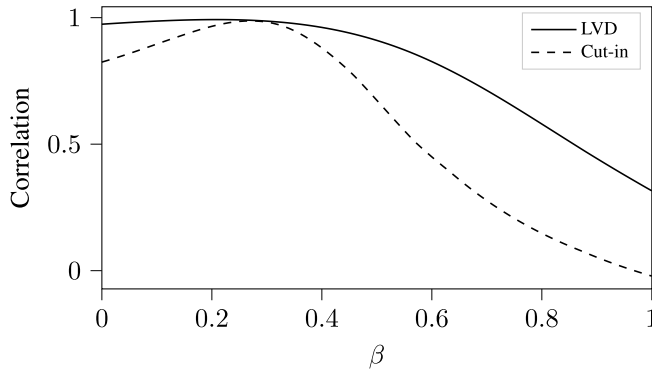


Fig. 11. Correlation between the medians of $M_1(\mathcal{W}^*, \mathcal{Z}^*, \mathcal{X}^*)$ (open circles in Figs. 9 and 10) and $\tilde{W}_1(\mathcal{Z}_{\text{large}}^*, \mathcal{W}^*)$ (filled circles in Figs. 9 and 10) for different values of β . The solid line shows the result for the LVD scenarios with a maximum correlation of 0.992 at $\beta \approx 0.21$. The dashed line shows the result for the cut-in scenarios with a maximum correlation of 0.987 at $\beta \approx 0.27$.

6) Repeat step 2.

7) Stop if $d_i = d_{i-1}$. Otherwise, return to step 3.

As an initial choice, $\beta_0 = 0.25$ seems appropriate. More specifically, when choosing $\beta_0 \in [0.1, 1]$, the optimal choice of d is found after one iteration.

VI. DISCUSSION

One of the advantages of our proposed method for generating scenario parameters is that less assumptions are needed regarding the parameterization of the scenarios:

- There is no assumption needed on a predetermined functional form of the time series data. For example, in an LVD scenario, the speed is often assumed to follow a polynomial function [10], a sinusoidal function, or a linear function [21]. In case of a predetermined functional form, parameters are fitted to the functional form. In our case, the SVD automatically determines the optimal choice of parameterization without relying on a predetermined functional form.
- There is no assumption needed for the shape of the distribution of the parameters. For example, a particular distribution, such as a Gaussian distribution [29] or a uniform distribution, may be assumed for which parameters are fitted. Alternatively, assumptions are made regarding the independence of the parameters [24]. In our case, the KDE automatically adapts its shape to the data and also considers the dependence among the different parameters.

It should be noted, however, that if there is a reason to believe that one or more of the assumptions are valid, than alternative methods for generation scenario parameters that make use of such assumptions might perform equally or better than the presented method [51]. In most cases, it will be difficult to provide a proper justification of the assumptions regarding the functional form of, e.g., the vehicle speed, and the pdf of the scenario parameters and the presented method will outperform methods relying on such assumptions. In any case, the presented SR metric provides an opportunity to verify the applicability of any assumptions regarding the scenario parameterization and parameter distributions.

The generated scenario parameters represent scenarios that could happen in real life and cover the same variety that is found in real-world traffic. Most likely, the majority of these scenarios are straightforward for the AV to deal with. To do an efficient assessment, the focus should be on scenarios that might lead to critical situations in which the probability of collision is high. That is why so-called *importance sampling* [52, Chapter 5.6] is often used for the assessment of AVs, e.g., see [5], [10], [53], [54]. With importance sampling, a different pdf, $g(\cdot)$, is used to sample scenario parameters, such that more emphasis is put on scenarios that might lead to critical situations. To get unbiased results, the result of a test with scenario parameters $\mathbf{x}$ is weighted by the ratio of the original probability density, $\hat{f}_H(\mathbf{x})$, and the probability density of the pdf used for importance sampling, $g(\mathbf{x})$ [10], [52]–[54]. Note that the importance sampling techniques explained in [10], [53], [54] can be directly applied on the estimated pdf $\hat{f}_H(\cdot)$ in (11) of the reduced set of parameters. In future work, our method for generating scenarios will be combined with importance sampling [10], [53], [54] for an assessment of an AV.

In some cases, one might want to sample from a conditional pdf, e.g., in case of sampling the scenario parameters for the LVD scenario such that the initial time gap equals a specified value. Sampling from a KDE such that one or more parameters are predetermined is straightforward [55]. In our case, sampling from $\hat{f}_h(\cdot)$ such that the time gap equals a specified value results in a linear constraint on the samples, because the reduced parameter vector $\tilde{\mathbf{v}}_i$ of (10) results from a linear mapping of the original parameters $\mathbf{x}_i$ of (2). In other words, one might want to sample v from $\hat{f}_H(\cdot)$ of (11), such that v is subject to the linear constraint

$$Av = b, \quad (25)$$

where A and b are a matrix and vector, respectively. In [40], an algorithm is provided for sampling from a pdf estimated using KDE such that the generated sample is subject to the constraint of (25). The main idea of [40] is to weight each parameter vector v_i , $i \in \{1, \dots, N_x\}$ in the KDE based on how closely the v_i matches the constraint (25).

The presented case study considers a vehicle for which the full trajectory is predetermined. For the presented scenarios, this works well, but the full trajectory is not predetermined in scenarios where the actor's behavior depends on the behavior of the ego vehicle [56]. To deal with such scenarios, one option is to use a driver behavior model (e.g., [57], [58]) with predefined parameters instead of describing the full trajectory. The parameters of the driver behavior model may be part of θ . The proposed method for generating scenario parameter values still applies in these kind of scenarios. Our ongoing research focuses on the assessment of AVs using scenarios in which driver behavior models are used for vehicles that may respond to the ego vehicle's behavior.

Since KDE is used, the generated scenario parameters represent variations of the data. Nevertheless, if the data do not contain scenarios that might lead to critical situations, such as an emergency braking maneuver or a reckless cut-in scenario, it is unlikely that such scenarios are generated,

even if importance sampling [10], [53], [54] is used. Therefore, when using the generated scenarios for the (safety) assessment of AVs, it is important that there is enough data such that the data contain such scenarios. Although there is no consensus yet on the required amount of data, some metrics have been proposed [39], [59] for determining whether enough data have been collected when using the data for the assessment of AVs.

This work employs the Wasserstein metric to propose the SR metric for evaluating the generated scenario parameters. It is illustrated how our proposed metric could be used to determine the appropriate number of parameters (d) and the type of distribution that is used to model the pdf of the scenario parameters. Also, the bandwidth h or bandwidth matrix H could also be determined by optimizing the proposed metric. In case of the bandwidth estimation, the disadvantage is that it would require more computational resources compared to, e.g., leave-one-out cross-validation.

More research is needed to determine the influences on the optimal choice for the penalty weight β . The case study has demonstrated one way to verify whether the initial choice of β was appropriate, but we do not yet know *why* a weight of $\beta \approx 0.25$ is an appropriate choice. The actual choice might depend on, among others, N_x , N_z , N_w , and the shape of the underlying distribution of the scenario parameters. Future research with a larger data set will allow us to better determine the optimal β and how this optimal value is influenced.

Future work involves researching the use of the proposed metric in combination with alternative methods for generating scenarios for the assessment of AVs. For example, Spooner *et al.* [28] have used a GAN [60] to create pedestrian crossing scenarios. One of the difficulties with GANs is to know when the GAN truly replicates the underlying distribution. Several metrics have been proposed [61] to evaluate the performance of GANs, among which a metric based on the Wasserstein metric that compares the generated data with test data. Alternatively, our proposed metric, which also considers the training data, could be considered for evaluating GANs. To judge the potential of our proposed metric in this application, more research is needed.

VII. CONCLUSION

It is essential for the deployment of Automated Vehicles (AVs) to develop assessment methods. Scenario-based assessment in which test cases are derived from real-world road traffic scenarios is regarded as a viable approach for assessing AVs. This work has presented a method to generate parameterized scenarios for the use in test case descriptions for the assessment of AVs. To not rely on a small set of parameters, we have used Singular Value Decomposition (SVD) to reduce the parameters. Parameter values for the scenarios are generated by drawing samples from the estimated probability density function (pdf) of the reduced set of parameters. To deal with the unknown shape of the pdf, it has been proposed to estimate the pdf using Kernel Density Estimation (KDE). This work has also presented a novel metric, the so-called Scenario

Representativeness (SR) metric, based on the Wasserstein metric, for evaluating whether the generated scenario parameters represent realistic scenarios while covering the same variety that is found in real-world traffic.

A case study has illustrated the proposed method for generating scenario parameter values using scenarios with a leading vehicle that decelerates and scenarios with a vehicle that performs a cut-in. The case study has also illustrated that the proposed metric correctly quantifies the degree to which the generated scenario parameter values represent real-world scenarios and, at the same time, cover the same variety of scenarios that is found in real life.

Future work involves applying the proposed method for more complex scenarios, e.g., scenarios that contain several different actors, to generate scenario-based test cases for the safety assessment of AVs. Additionally, it would be of interest to apply importance sampling for AV assessment in combination with the proposed method for generating scenarios. Other future work involves investigating the use of the proposed metric in combination with alternative methods for generating scenarios for the assessment of AVs.

REFERENCES

- [1] K. Bengler, K. Dietmayer, B. Farber, M. Maurer, C. Stiller, and H. Winner, "Three decades of driver assistance systems: Review and future perspectives," *IEEE Intell. Transp. Syst. Mag.*, vol. 6, no. 4, pp. 6–22, Oct. 2014.
- [2] J. E. Stellet, M. R. Zofka, J. Schumacher, T. Schamm, F. Niewels, and J. M. Zöllner, "Testing of advanced driver assistance towards automated driving: A survey and taxonomy on existing approaches and open questions," in *Proc. 18th Int. Conf. Intell. Transp. Syst.*, Sep. 2015, pp. 1455–1462.
- [3] P. Koopman and M. Wagner, "Challenges in autonomous vehicle testing and validation," *SAE Int. J. Transp. Saf.*, vol. 4, no. 1, pp. 15–24, Apr. 2016.
- [4] N. Kalra and S. M. Paddock, "Driving to safety: How many miles of driving would it take to demonstrate autonomous vehicle reliability?" *Transp. Res. A, Policy Pract.*, vol. 94, pp. 182–193, Dec. 2016.
- [5] D. Zhao, X. Huang, H. Peng, H. Lam, and D. J. LeBlanc, "Accelerated evaluation of automated vehicles in car-following maneuvers," *IEEE Trans. Intell. Transp. Syst.*, vol. 19, no. 3, pp. 733–744, Mar. 2018.
- [6] S. Riedmaier, T. Ponn, D. Ludwig, B. Schick, and F. Diermeyer, "Survey on scenario-based safety assessment of automated vehicles," *IEEE Access*, vol. 8, pp. 87456–87477, 2020.
- [7] H. Elrofai, J.-P. Paardekooper, E. de Gelder, S. Kalisvaart, and O. Op den Camp, "Scenario-based safety validation of connected and automated driving," Netherlands Org. Appl. Sci. Res., TNO, The Hague, The Netherlands, Tech. Rep., 2018. [Online]. Available: <http://publications.tno.nl/publication/34626550/AyT8Zc/TNO-2018-streetwise.pdf>
- [8] A. Pütz, A. Zlocki, J. Bock, and L. Eckstein, "System validation of highly automated vehicles with a database of relevant traffic scenarios," in *Proc. 12th ITS Eur. Congr.*, 2017, pp. 1–8. [Online]. Available: https://www.pegasusprojekt.de/files/tmpl/pdf/12th%20ITS%20European%20Co%20ngress_Folien.pdf
- [9] R. Krajewski, J. Bock, L. Kloecker, and L. Eckstein, "The high D dataset: A drone dataset of naturalistic vehicle trajectories on German highways for validation of highly automated driving systems," in *Proc. 21st Int. Conf. Intell. Transp. Syst. (ITSC)*, Nov. 2018, pp. 2118–2125.
- [10] E. de Gelder and J.-P. Paardekooper, "Assessment of automated driving systems using real-life scenarios," in *Proc. Intell. Veh. Symp. (IV)*, 2017, pp. 589–594.
- [11] J. Antona-Makoshi, N. Uchida, K. Yamazaki, K. Ozawa, E. Kitahara, and S. Taniguchi, "Development of a safety assurance process for autonomous vehicles in Japan," in *Proc. 26th Int. Tech. Conf. Enhanced Saf. Veh. (ESV)*, 2019, pp. 1–18. [Online]. Available: <https://www-esv.nhtsa.dot.gov/Proceedings/26/26ESV-000286.pdf>
- [12] G. H. Golub and C. F. Van Loan, *Matrix Computations*, vol. 3. Baltimore, MD, USA: John Hopkins Univ. Press, 2013.

- [13] M. Rosenblatt, "Remarks on some nonparametric estimates of a density function," *Ann. Math. Statist.*, vol. 27, no. 3, pp. 832–837, 1956.
- [14] E. Parzen, "On estimation of a probability density function and mode," *Ann. Math. Statist.*, vol. 33, no. 3, pp. 1065–1076, Sep. 1962.
- [15] L. Rüschendorf, "The Wasserstein distance and approximation theorems," *Probab. Theory Rel. Fields*, vol. 70, no. 1, pp. 117–129, 1985.
- [16] L. Li, W.-L. Huang, Y. Liu, N.-N. Zheng, and F.-Y. Wang, "Intelligence testing for autonomous vehicles: A new approach," *IEEE Trans. Intell. Veh.*, vol. 1, no. 2, pp. 158–166, Jun. 2016.
- [17] U. Lages, M. Spencer, and R. Katz, "Automatic scenario generation based on laserscanner reference data and advanced offline processing," in *Proc. Intell. Vehicles Symp. (IV)*, Jun. 2013, pp. 146–148.
- [18] M. R. Zofka, F. Kuhnt, R. Kohlhaas, C. Rist, T. Schamm, and J. M. Zöllner, "Data-driven simulation and parametrization of traffic scenarios for the development of advanced driver assistance systems," in *Proc. 18th Int. Conf. Inf. Fusion*, Jul. 2015, pp. 1422–1428. [Online]. Available: <https://ieeexplore.ieee.org/abstract/document/7266724>
- [19] L. Stepien *et al.*, "Applying heuristics to generate test cases for automated driving safety evaluation," *Appl. Sci.*, vol. 11, no. 21, p. 10166, Oct. 2021.
- [20] S. Feng, Y. Feng, C. Yu, Y. Zhang, and H. X. Liu, "Testing scenario library generation for connected and automated vehicles, Part I: Methodology," *IEEE Trans. Intell. Transp. Syst.*, vol. 22, no. 3, pp. 1573–1582, Mar. 2021.
- [21] S. Thal *et al.*, "Incorporating safety relevance and realistic parameter combinations in test-case generation for automated driving safety assessment," in *Proc. Intell. Transp. Syst. Conf. (ITSC)*, 2020, pp. 666–671.
- [22] S. Feng, Y. Feng, X. Yan, S. Shen, S. Xu, and H. X. Liu, "Safety assessment of highly automated driving systems in test tracks: A new framework," *Accident Anal. Prevention*, vol. 144, Mar. 2020, Art. no. 105664.
- [23] L. Li, N. Zheng, and F.-Y. Wang, "A theoretical foundation of intelligence testing and its application for intelligent vehicles," *IEEE Trans. Intell. Transp. Syst.*, vol. 22, pp. 6297–6306, 2020.
- [24] S. Feng, X. Yan, H. Sun, Y. Feng, and H. X. Liu, "Intelligent driving intelligence test for autonomous vehicles with naturalistic and adversarial environment," *Nature Commun.*, vol. 12, no. 1, pp. 1–14, Dec. 2021.
- [25] M. Koren, S. Alsaif, R. Lee, and M. J. Kochenderfer, "Adaptive stress testing for autonomous vehicles," in *Proc. Intell. Veh. Symp. (IV)*, 2018, pp. 1898–1904.
- [26] A. Corso and M. J. Kochenderfer, "Interpretable safety validation for autonomous vehicles," in *Proc. 23rd Int. Conf. Intell. Transp. Syst. (ITSC)*, 2020, pp. 1–6.
- [27] F. Schuldt, A. Reschka, and M. Maurer, "A method for an efficient, systematic test case generation for advanced driver assistance systems in virtual environments," in *Automotive Systems Engineering II*. Cham, Switzerland: Springer, 2018, pp. 147–175.
- [28] J. Spooner, V. Palade, M. Cheah, S. Kanarachos, and A. Daneshkhan, "Generation of pedestrian crossing scenarios using ped-cross generative adversarial network," *Appl. Sci.*, vol. 11, no. 2, p. 471, Jan. 2021.
- [29] O. Gietelink, "Design and validation of advanced driver assistance systems," Ph.D. dissertation, Transp. Res. Centre Delft, Delft Univ. Technol., Delft, The Netherlands, 2007. [Online]. Available: <http://resolver.tudelft.nl/uuid:b2f0e7f6-6255-4932-8b5e-d3ef67cd81ec>
- [30] S. H. Cha, "Comprehensive survey on distance/similarity measures between probability density functions," *Int. J. Math. Models Methods Appl. Sci.*, vol. 1, no. 4, pp. 300–307, Nov. 2007.
- [31] S. Kullback and R. A. Leibler, "On information and sufficiency," *Ann. Math. Statist.*, vol. 22, no. 1, pp. 79–86, 1951.
- [32] D. Zhao *et al.*, "Accelerated evaluation of automated vehicles safety in lane-change scenarios based on importance sampling techniques," *IEEE Trans. Intell. Transp. Syst.*, vol. 18, no. 3, pp. 595–607, Mar. 2017.
- [33] D. W. Scott, *Multivariate Density Estimation: Theory, Practice, and Visualization*. Hoboken, NJ, USA: Wiley, 1992.
- [34] E. de Gelder *et al.*, "Towards an ontology for scenario definition for the assessment of automated vehicles: An object-oriented framework," *IEEE Trans. Intell. Veh.*, early access, Jan. 25, 2022, doi: [10.1109/TIV.2022.3144803](https://doi.org/10.1109/TIV.2022.3144803).
- [35] C. de Boor, *A Practical Guide to Splines*, vol. 27. New York, NY, USA: Springer, 1978.
- [36] H. Abdi and L. J. Williams, "Principal component analysis," *Wiley Interdiscipl. Rev., Comput. Statist.*, vol. 2, no. 4, pp. 433–459, 2010.
- [37] B. A. Turlach, "Bandwidth selection in kernel density estimation: A review," Inst. Statistik Ökonometrie, Humboldt-Univ. Berlin, Berlin, Germany, Tech. Rep., 1993.
- [38] T. Duong, "KS: Kernel density estimation and kernel discriminant analysis for multivariate data in R," *J. Stat. Softw.*, vol. 21, no. 7, pp. 1–16, 2007.
- [39] E. de Gelder, J.-P. Paardekooper, O. Op den Camp, and B. De Schutter, "Safety assessment of automated vehicles: How to determine whether we have collected enough field data?" *Traffic Injury Prevention*, vol. 20, no. S1, pp. 162–170, 2019.
- [40] E. de Gelder, E. Cator, J.-P. Paardekooper, O. Op den Camp, and B. De Schutter, "Constrained sampling from a kernel density estimator to generate scenarios for the assessment of automated vehicles," in *Proc. Intell. Veh. Symp. Workshops*, 2021, pp. 203–208.
- [41] R. P. W. Duin, "On the choice of smoothing parameters for Parzen estimators of probability density functions," *IEEE Trans. Comput.*, vol. C-25, no. 11, pp. 1175–1179, Nov. 1976.
- [42] A. Z. Zambom and R. Dias, "A review of kernel density estimation with applications to econometrics," *Int. Econ. Rev.*, vol. 5, no. 1, pp. 20–42, 2013. [Online]. Available: <http://www.era.org.tr/makaleler/13120083.pdf>
- [43] A. Gramacki, *Nonparametric Kernel Density Estimation Its Comput. Aspects*, J. Kacprzyk, Ed. Cham, Switzerland: Springer, 2018.
- [44] Y. Rubner, C. Tomasi, and L. J. Guibas, "The earth mover's distance as a metric for image retrieval," *Int. J. Comput. Vis.*, vol. 40, no. 2, pp. 99–121, Nov. 2000.
- [45] M. Sommerfeld and A. Munk, "Inference for empirical Wasserstein distances on finite spaces," *J. Roy. Stat. Soc., Ser. B Stat. Methodol.*, vol. 80, no. 1, pp. 219–238, Jan. 2018.
- [46] J.-P. Paardekooper *et al.*, "Automatic identification of critical scenarios in a public dataset of 6000 km of public-road driving," in *26th Int. Tech. Conf. Enhanced Saf. Veh. (ESV)*, 2019, pp. 1–8. [Online]. Available: <https://www-esv.nhtsa.dot.gov/Proceedings/26/26ESV-000255.pdf>
- [47] J. Elfving, R. Appeldoorn, S. van den Dries, and M. Kwakernaat, "Effective world modeling: Multisensor data fusion methodology for automated driving," *Sensors*, vol. 16, no. 10, pp. 1–27, 2016.
- [48] E. de Gelder, J. Manders, C. Grappiolo, J.-P. Paardekooper, O. Op den Camp, and B. De Schutter, "Real-world scenario mining for the assessment of automated vehicles," in *Proc. Int. Transp. Syst. Conf. (ITSC)*, 2020, pp. 1073–1080.
- [49] B. Doerr and A. M. Sutton, "When resampling to cope with noise, use median, not mean," in *Proc. Genetic Evol. Comput. Conf.*, Jul. 2019, pp. 242–248.
- [50] B. Efron, "Bootstrap methods: Another look at the jackknife," in *Breakthroughs Statistics*. New York, NY, USA: Springer, 1992, pp. 569–593.
- [51] S. Siegel, "Nonparametric statistics," *Amer. Stat.*, vol. 11, no. 3, pp. 13–19, Jun. 1957.
- [52] R. Y. Rubinstein and D. P. Kroese, *Simulation Monte Carlo Method*. Hoboken, NJ, USA: Wiley, 2016.
- [53] S. Jesenski, N. Tiemann, J. E. Stellet, and J. M. Zöllner, "Scalable generation of statistical evidence for the safety of automated vehicles by the use of importance sampling," in *Proc. 23rd Int. Conf. Intell. Transp. Syst. (ITSC)*, 2020, pp. 1–8.
- [54] Y. Xu, Y. Zou, and J. Sun, "Accelerated testing for automated vehicles safety evaluation in cut-in scenarios based on importance sampling, genetic algorithm and simulation applications," *J. Intell. Connected Veh.*, vol. 1, no. 1, pp. 28–38, Oct. 2018.
- [55] M. P. Holmes, A. G. Gray, and C. Lee Isbell, "Fast nonparametric conditional density estimation," 2012, *arXiv:1206.5278*.
- [56] M. Althoff, M. Koschi, and S. Manzing, "CommonRoad: Composable benchmarks for motion planning on roads," in *Proc. IEEE Intell. Veh. Symp. (IV)*, Jun. 2017, pp. 719–726.
- [57] M. Treiber, A. Hennecke, and D. Helbing, "Congested traffic states in empirical observations and microscopic simulations," *Phys. Rev. E, Stat. Phys. Plasmas Fluids Relat. Interdiscip. Top.*, vol. 62, no. 2, pp. 1805–1824, Aug. 2000.
- [58] A. Kesting, M. Treiber, and D. Helbing, "General lane-changing model MOBIL for car-following models," *Transp. Res. Rec.*, vol. 1999, pp. 86–94, Jan. 2007.
- [59] W. Wang, C. Liu, and D. Zhao, "How much data are enough? A statistical approach with case study on longitudinal driving behavior," *IEEE Trans. Intell. Veh.*, vol. 2, no. 2, pp. 85–98, Jun. 2017.

- [60] I. J. Goodfellow *et al.*, "Generative adversarial nets," in *Proc. 27th Int. Conf. Neural Inf. Process. Syst. (NIPS)*, vol. 2, 2014, pp. 2672–2680.
- [61] A. Borji, "Pros and cons of GAN evaluation measures," *Comput. Vis. Image Understand.*, vol. 179, pp. 41–65, Feb. 2019.



Erwin de Gelder received the M.Sc. degree (*cum laude*) in systems and control from the Delft University of Technology, Delft, The Netherlands, in 2014, where he is currently pursuing the Ph.D. degree on the topic of assessment procedures for automated vehicles. Since 2014, he has been with the Netherlands Organization for Applied Scientific Research (TNO). From 2017 to 2019, he worked with Nanyang Technological University, Singapore, to apply his research for an assessment procedure for automated vehicles in Singapore. His current

research focuses on a quantitative assessment methodology for automated vehicles using scenarios obtained from real-world data.



Jasper Hof received the M.Sc. degree (*cum laude*) in mathematics from Radboud University in 2020, where he specialized in applied stochastics. He is currently pursuing the Ph.D. degree in statistical genetics with the Radboud University Medical Center. During his M.Sc., he researched methods that enable the generation of scenarios for assessment of automated vehicles at the Netherlands Organization for Applied Scientific Research (TNO). During his Ph.D., he develops statistical methods that aim to statistically infer associations between

recurrent events and genetic variants in genome-wide association studies.



Eric Cator is currently a Professor in applied stochastic and the Director of the IMAPP Institute, Radboud University. His research interests are in probability and statistics, focusing on proving asymptotic properties of stochastic models, such as the contact process or specific estimators in non-parametric statistical models. He is also actively involved in interdisciplinary collaborations with astronomy and biology, developing state-of-the-art statistical methods for data specific to these fields, and he has given advice to many companies and

research groups on their data analysis problems.



Jan-Pieter Paardekooper received the M.Sc. degree in astrophysics from Leiden University, Leiden, The Netherlands, in 2006, and the Ph.D. degree in theoretical astrophysics in Leiden, in 2010. He is currently a Researcher with TNO, The Netherlands, on the topic of artificial intelligence (AI) in connected and automated vehicles. Since 2018, he has been a part-time Research Fellow with the Department of Artificial Intelligence, Donders Institute for Brain, Cognition and Behaviour, Radboud University, Nijmegen, The Netherlands.

His research interests include AI safety, neuro-symbolic AI, and situation awareness for AI.



Olaf Op den Camp received the M.Sc. degree (*cum laude*) in mechanical engineering and M.Sc. degree (*cum laude*) in biomedical engineering from the Eindhoven University of Technology, and the Ph.D. degree with a study into the identification of material parameters of biological tissue in 1996. In 1995, he joined the Netherlands Organization for Applied Scientific Research (TNO). He is a Senior Consultant and the Technical Lead of the StreetWise scenario database development. Since 2018, he has been working part time with Nanyang Technological

University, Singapore, to develop and implement a scenario-based safety assessment framework to support safe deployment of autonomous vehicles within the CETRAN Program.



Jeroen Ploeg received the M.Sc. degree in mechanical engineering from the Delft University of Technology, Delft, The Netherlands, in 1988, and the Ph.D. degree in mechanical engineering on the control of vehicle platoons from the Eindhoven University of Technology, Eindhoven, The Netherlands, in 2014.

He is currently with 2getthere, Utrecht, The Netherlands, where he leads the research activities in the field of cooperative automated driving for automated transit systems. Since 2017, he has been

holding the position of a part-time Associate Professor with the Mechanical Engineering Department, Eindhoven University of Technology. His research interests include control system design for cooperative and automated vehicles.



Bart de Schutter (Fellow, IEEE) received the Ph.D. degree (*summa cum laude*) in applied sciences from KU Leuven, Belgium, in 1996.

He is currently a Full Professor and the Head of the Department with the Delft Center for Systems and Control, Delft University of Technology, Delft, The Netherlands. His current research interests include multi-level and multi-agent control, learning-based control, control of hybrid systems with applications in intelligent transportation systems, and smart energy systems. He is also a Senior Editor of the

IEEE TRANSACTIONS ON INTELLIGENT TRANSPORTATION SYSTEMS and an Associate Editor of IEEE TRANSACTIONS ON AUTOMATIC CONTROL.