



IDSplat: Instance-Decomposed 3D Gaussian Splatting for Driving Scenes

Downloaded from: <https://research.chalmers.se>, 2026-08-28 16:20 UTC

Citation for the original published paper (version of record):

Lindström, C., Rafidashti, M., Fatemi, M. et al (2026). IDSplat: Instance-Decomposed 3D Gaussian Splatting for Driving Scenes. Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition: 316-326

N.B. When citing this work, cite the original published paper.

© 2026 IEEE. Personal use of this material is permitted. Permission from IEEE must be obtained for all other uses, in any current or future media, including reprinting/republishing this material for advertising or promotional purposes, or reuse of any copyrighted component of this work in other works.

IDSplat: Instance-Decomposed 3D Gaussian Splatting for Driving Scenes

Carl Lindström^{†,1,2} Mahan Rafidashti^{†,1,2} Maryam Fatemi¹
 Lars Hammarstrand² Martin R. Oswald³ Lennart Svensson²

¹Zenseact ²Chalmers University of Technology ³University of Amsterdam
 {firstname.lastname}@{zenseact.com, chalmers.se}

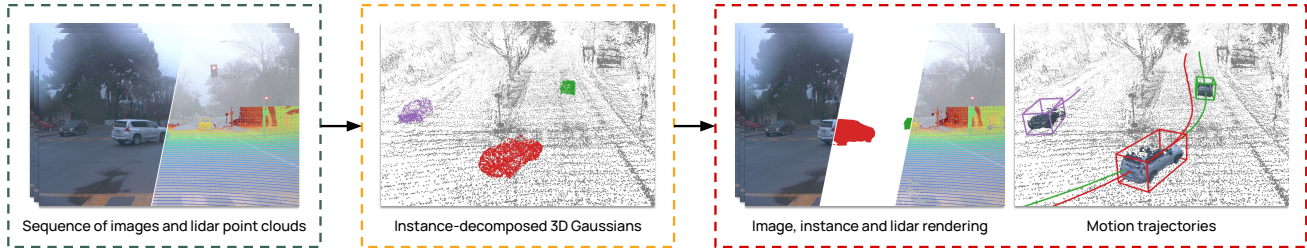


Figure 1. IDSplat performs self-supervised reconstruction of dynamic scenes with explicit instance-decomposition and learnable motion trajectories. IDSplat enables high-fidelity rendering of images, instances, and lidar point clouds without the need for human annotations.

Abstract

Reconstructing dynamic driving scenes is essential for developing autonomous systems through sensor-realistic simulation. Although recent methods achieve high-fidelity reconstructions, they either rely on costly human annotations for object trajectories or use time-varying representations without explicit object-level decomposition, leading to intertwined static and dynamic elements that hinder scene separation. We present IDSplat, a self-supervised 3D Gaussian Splatting framework that reconstructs dynamic scenes with explicit instance decomposition and learnable motion trajectories, without requiring human annotations. Our key insight is to model dynamic objects as coherent instances undergoing rigid transformations, rather than unstructured time-varying primitives. For instance decomposition, we employ zero-shot, language-grounded video tracking anchored to 3D using lidar, and estimate consistent poses via feature correspondences. We introduce a coordinated-turn smoothing scheme to obtain temporally and physically consistent motion trajectories, mitigating pose misalignments and tracking failures, followed by joint optimization of object poses and Gaussian parameters. Experiments on the Waymo Open Dataset demonstrate that our method achieves competitive reconstruction quality while maintaining instance-level decomposition and generalizes across diverse sequences and view densities without retraining, making it practical for large-scale autonomous driving applications. Code will be released.

1. Introduction

Reconstructing dynamic scenes has become an important cornerstone in the development of autonomous driving systems, enabling closed-loop training and testing through sensor-realistic renderings. In contrast to real-world testing, digital twins provide a scalable, low-cost, and safe means of exploring novel driving scenarios using already collected data [16, 17]. Recent work has improved the quality, efficiency, and sensor compatibility of such reconstructions [4, 8, 30, 31, 36, 39, 42, 48], but typically rely on human-annotated object trajectories and 3D bounding boxes, which are expensive and time-consuming to acquire at scale.

To address this challenge, several self-supervised approaches have emerged that aim to reconstruct dynamic scenes without human annotations [3, 9, 10, 20, 40, 46]. Although they produce high-quality renderings, they lack explicit instance decomposition, significantly limiting their practical use, since novel scenario generation requires manipulating individual dynamic objects.

In this paper, we address the problem of reconstructing realistic dynamic scenes without human annotations, while preserving instance-level decomposition and learning the underlying motion trajectories of individual objects. Reconstructing dynamic scenes using 3D Gaussian Splatting (3DGS) [11] presents a significant challenge, as it violates the assumption that each 3D point maintains a fixed position and appearance across viewpoints. While prior self-supervised approaches address this by introducing time-dependent Gaussian parameters [3, 9, 20], we instead preserve the geometry and appearance of coherent object in-

[†]These authors contributed equally to this work.

stances over time and optimize their rigid transformations to capture the true underlying motion. This formulation preserves instance-level decomposition and enables controllable trajectories, rather than relying on time-varying primitives whose changing visibility and appearance can lead to inconsistent scene representations.

Obtaining instance-level decomposition of 3D Gaussians without annotated poses introduces an additional challenge. Although 3D object trackers can perform well in estimating object poses and have been used effectively in dynamic scene reconstruction [30, 39], they rely on human annotations for fine-tuning on the target data and are constrained by the predefined taxonomy of those annotations. To overcome this limitation, we leverage recent advances in vision models and employ a language-grounded video tracker to extract instance masks in a zero-shot manner, allowing both generalization to new datasets and new classes without retraining or additional annotations. These masks are lifted to 3D using corresponding lidar point clouds, and object poses are estimated via RANSAC using DINOv3 [24] feature correspondences. To further address pose misalignments and missing detections, we introduce an iterative robust coordinated-turn smoothing scheme that discards outliers and refines the trajectories.

Although accurate initialization of object trajectories is crucial, inaccuracies are inevitable due to missing instance masks, inaccurate lidar poses, or tracking drift over time. To mitigate these errors, we make a final refinement of the object trajectories guided by reconstruction errors during the 3D Gaussian Splatting optimization.

We present IDSplat, a novel 3D Gaussian Splatting framework designed to handle dynamic scenes through instance-level decomposition and joint optimization of object appearance, geometry, and motion. We extensively evaluate the effectiveness of our method on Waymo Open Dataset [27], achieving state-of-the-art results across a diverse set of sequences and test protocols. In summary, our contributions are as follows:

- We propose a self-supervised framework for dynamic scene reconstruction that explicitly decomposes scenes into object instances with learnable motion trajectories, enabling joint rendering, segmentation and motion tracking.
- We introduce a zero-shot approach for 3D instance decomposition and pose estimation, enabling generalization to new datasets and object classes without retraining or human annotations.
- We present simple yet effective techniques for optimizing and refining motion trajectories, combining motion modeling and photometric consistency to obtain accurate trajectories even under sparse views.

2. Related work

Annotation-based rendering for autonomous driving:

Neural radiance fields (NeRF) [19] inspired numerous NeRF-based methods for dynamic road scenes [4, 21, 30, 36, 42]. These achieve high-quality renderings and naturally extend to new sensors [21, 30], but are sample-intensive, resulting in slow rendering speeds and limiting their scalability.

3D Gaussian Splatting (3DGS) [11] provides an explicit rasterization-friendly representation, enabling orders-of-magnitude faster rendering. Automotive scene reconstruction with 3DGS has been explored in several works, including lidar-based extensions [2, 4, 8, 12, 39, 47, 48].

To model dynamic scenes, both NeRF- and 3DGS-based approaches typically rely on accurate 3D bounding boxes with temporal instance associations. These enable the scene to be decomposed into a static background and dynamic foreground components, with each dynamic object transformed according to its trajectory. To achieve high-fidelity reconstruction, state-of-the-art approaches typically rely on either human-annotated bounding boxes or predictions from high-performing 3D object trackers that have been carefully adapted to the target dataset. This reliance on curated annotations or dataset-specific trackers limits the scalability and zero-shot generalization of these methods to new datasets.

Self-supervised dynamic scene reconstruction: Self-supervised approaches also separate static and dynamic regions, using either separate hash-grids [31, 40], or time-varying 3DGS representations [3, 9, 18, 20, 26, 28, 34, 38, 41], where Gaussian attributes are allowed to change over time. Instead of bounding boxes, these methods rely on photometric and geometric consistency or foundation model outputs to guide decomposition, either via learned features or explicitly via predicted masks [9, 10, 20, 34, 40, 43, 46]. DeSiRe-GS [20] segments dynamic regions in images using features from FiT3D [44] and uses these dynamic masks to optimize time-varying Gaussians. Similarly, AD-GS [10] relies on Grounded-SAM-2 [23] masks, but uses B-splines and trigonometric functions to enforce smoother trajectories for individual Gaussians. Other works attempt more detailed decomposition, such as CoDa-4DGS [26] which enables semantic segmentation by learning per-Gaussian feature vectors. A key limitation of these models is that temporal changes are modeled at the primitive level, preventing decomposition into coherent object instances. This restricts practical applications such as novel scenario generation, auto-labeling, and simulation.

IDSplat addresses this by decomposing scenes into a static background and dynamic foreground instances that remain consistent across the sequence. Each instance’s motion trajectory is explicitly represented, enabling re-assignment, manipulation or removal. We achieve self-supervised instance decomposition using Grounded-SAM-2

for zero-shot masks and DINOv3 [24] for robust registration across time, and guide the refinement of motion trajectories through photometric and geometric consistency.

3. Method

Given a sequence of images and lidar point clouds from a driving scenario, our goal is to learn a 3D representation of the scene, including instance-decomposed dynamic objects with associated learnable motion trajectories. In the following, we describe our scene representation (Sec. 3.1), how the representation is decomposed into separate instances (Sec. 3.2), and how associated trajectories are estimated and refined (Secs. 3.3 and 3.4). Finally, we describe how the complete scene is optimized (Sec. 3.5). See Fig. 2 for an overview of our method.

3.1. Scene representation

We represent the scene by a set of translucent 3D Gaussians, parameterized with occupancy probability $o \in [0, 1]$, mean $\mu \in \mathbb{R}^3$, and covariance $\Sigma \in \mathbb{R}^{3 \times 3}$. Together, the parameters describe the position, extent and visibility of the Gaussian. To facilitate both camera and lidar rendering, we follow SplatAD [8] and assign each Gaussian a feature vector $\mathbf{f}^{\text{rgb}} \in \mathbb{R}^3$ to represent its base color, and another feature vector $\mathbf{f} \in \mathbb{R}^{D_f}$ to represent view-dependent effects and lidar properties. Further, each Gaussian also has an associated discrete ID denoted by $z \in \{0, \dots, N_{\text{ID}}\}$, determining which instance it belongs to, or whether it belongs to the static background ($z = 0$). We also adopt sensor-specific embeddings to model appearance shifts.

To account for scene dynamics, we parameterize the motion of Gaussians using SE(3) poses. We assume each instance corresponds to a rigid object and transform all Gaussians associated with a given ID using the same rigid transformation. The position of Gaussian i in the world at time t is given as

$$\mu_{i,t} = \mathbf{T}_{z_i,t} \mu_i \quad (1)$$

where $\mu_{i,t}$ denotes the Gaussian’s position in world coordinates at time t , μ_i is its position in the canonical coordinate system of instance z_i , and $\mathbf{T}_{z_i,t} \in \text{SE}(3)$ represents the transformation from that canonical frame to the world at time t . Gaussians associated with the static background are defined directly in the world coordinate system and are not transformed.

3.2. Instance decomposition

To obtain object instances, we employ Grounded-SAM-2 [23] to generate instance masks from video frames using class prompts. As our scene is represented in 3D, we lift the 2D instance information to 3D by projecting the temporally closest lidar points onto the image plane and assigning them the corresponding instance IDs. To reduce outliers from inaccurate projections, arising from line-of-sight mismatches

due to camera-lidar mounting offsets or large uncorrected object motions, we apply a two-stage filtering process. We first erode the instance masks prior to projection to reduce the influence from sensor mounting offsets, and then cluster the projected points using DBSCAN [6], retaining only the largest cluster for each instance.

3.3. Trajectory estimation

With the instance points identified, we estimate the trajectory of each object instance by registering the lidar points associated with that instance. For an instance z , we first define a canonical frame centered at the midpoint of the axis-aligned bounding box enclosing the points from the time step with the densest lidar observations. This defines the initial pose as

$$\mathbf{T}_{z,t_{\text{init}}(z)} = \begin{bmatrix} \mathbf{I}_{3 \times 3} & \mathbf{c}_z \\ \mathbf{0}^T & 1 \end{bmatrix}, \quad (2)$$

where $t_{\text{init}}(z)$ denotes the time with the highest lidar density for instance z , and $\mathbf{c}_z$ is the center of the corresponding bounding box. Subsequent frames are registered to the canonical frame, $\mathbf{T}_{z,t_{\text{init}}(z)}$, ordered by point density. We extract DINOv3 [24] features from image projections and establish correspondences between frames based on cosine similarity. The rigid transformation is then estimated using RANSAC, where each hypothesis is computed via the Umeyama estimator [32] from three randomly sampled correspondences. The pose with the largest number of structural inliers is selected if the registration is deemed successful. We consider the registration successful when the inlier ratio of the target point cloud exceeds a predefined threshold, motivated by the fact that the target (from frame t_j) will always be smaller than the canonical source. When a registration succeeds, the resulting pose for time t_j is added to the trajectory and given by

$$\mathbf{T}_{z,t_j} = \mathbf{T}_{z,t_j \leftarrow t_{\text{init}}(z)} \mathbf{T}_{z,t_{\text{init}}(z)}, \quad (3)$$

where $\mathbf{T}_{z,t_j \leftarrow t_{\text{init}}(z)}$ is the rigid transformation estimated from RANSAC, and $\mathbf{T}_{z,t_j}$ represents the pose of instance z at time t_j . The corresponding points from frame t_j are then transformed into the canonical frame using the inverse of $\mathbf{T}_{z,t_j \leftarrow t_{\text{init}}(z)}$ and merged into the canonical point set.

3.4. Trajectory smoothing

While the RANSAC-based registration yields initial pose estimates between pairs of point clouds, our goal is to estimate a temporally and physically consistent object trajectory, and reduce the impact from imperfections due to registration errors or temporal gaps from missing instance masks. To address this, we refine the trajectories through an iterative coordinated-turn (CT) smoothing formulated as a pose graph optimization problem.

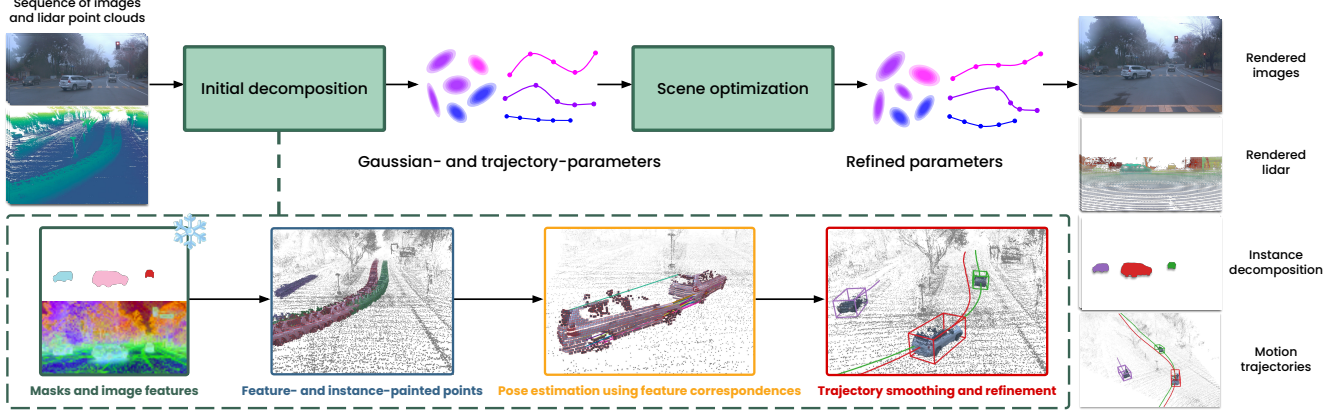


Figure 2. Overview of our method. 2D masks from Grounded-SAM-2 are lifted to 3D using corresponding lidar point clouds to initialize instances. Object poses are estimated via RANSAC using DINOv3 feature correspondences and further refined through iterative CT smoothing. Trajectories and Gaussian parameters are then optimized to render images, lidar, and instances with motion trajectories.

Optimization formulation: We employ GTSAM [5] to optimize a state vector comprising poses $\mathbf{T}_t \in \text{SE}(3)$, translational speeds $v_t \in \mathbb{R}$, and curvatures $\kappa_t \in \mathbb{R}$ for all timesteps t in the sequence. The estimated poses from RANSAC are added as noisy measurement factors through Gaussian likelihoods with a Huber loss function. Since instance masks may be intermittent, measurements may only be available at a subset of timesteps. We incorporate CT motion model factors to encourage physically grounded, temporally smooth trajectory states and to reject measurement outliers. To align the axis-aligned measurements with the motion model, we also estimate a single rotation, shared across all times, that orients the local x-axis along the direction of motion. Smoothness priors on the trajectory is further added as random walk processes on speed and curvature, and additional priors are added to encourage small roll and pitch angles as well as moderate curvature values.

Outlier rejection: To address outliers, we apply iterative refinement of the trajectories by first performing a single optimization pass and then identifying and removing measurements whose residuals exceed a predefined threshold. The optimization is then re-run with the pruned measurement set, improving the robustness of trajectory estimates even in the presence of registration failures.

3.5. Scene optimization

Our scene representation consists of a set of static Gaussians and a set of dynamic Gaussians associated to instances with corresponding trajectories. All Gaussian parameters and trajectories are optimized jointly through self-supervised reconstruction of images and lidar point clouds. We adopt the rasterization proposed in [8] and optimize the entire model jointly using the reconstruction loss

$$\mathcal{L} = \lambda_r \mathcal{L}_1 + (1 - \lambda_r) \mathcal{L}_{\text{SSIM}} + \mathcal{L}_{\text{lidar}} + \lambda_{\text{MCMC}} \mathcal{L}_{\text{MCMC}}, \quad (4)$$

where $\mathcal{L}_1$ and $\mathcal{L}_{\text{SSIM}}$ are L1 and SSIM losses on the rendered images, and $\mathcal{L}_{\text{MCMC}}$ denotes the opacity and scale regularization used in [13]. $\mathcal{L}_{\text{lidar}}$ is the loss from lidar reconstruction and is defined as

$$\mathcal{L}_{\text{lidar}} = \lambda_{\text{depth}} \mathcal{L}_{\text{depth}} + \lambda_{\text{los}} \mathcal{L}_{\text{los}} + \lambda_{\text{inten}} \mathcal{L}_{\text{inten}} + \lambda_{\text{raydrop}} \mathcal{L}_{\text{BCE}}, \quad (5)$$

where $\mathcal{L}_{\text{depth}}$ and $\mathcal{L}_{\text{inten}}$ are L2 losses on the rendered expected lidar range and intensity, and $\mathcal{L}_{\text{los}}$ is a line-of-sight loss that penalize accumulated opacity before the ground truth lidar range. $\mathcal{L}_{\text{BCE}}$ is a binary cross-entropy loss on predicted ray drop probability. For ease of comparison, we adopt the hyperparameters from [8]. See Sec. A for details.

Dynamic Gaussians are initialized from the canonical point sets created during point registration (3.3), while static Gaussians are seeded from lidar points not associated with any instance. RGB values for both sets are assigned by projecting the corresponding lidar points into the temporally closest image. Additional Gaussians are sampled randomly within the lidar range, and sampled linearly in disparity beyond observed points. Following [8], we employ the MCMC densification strategy from [13].

4. Experiments

To thoroughly evaluate IDSplat, we benchmark its performance on novel view synthesis (NVS) and image reconstruction using Waymo Open Dataset [27]. We compare against state-of-the-art self-supervised methods adapted to automotive scenes under multiple evaluation protocols, spanning a range of view densities and dynamic object categories. To assess generalization, we conduct additional ex-

Table 1. NVS results on experimental settings of AD-GS and DeSiRe-GS. Results for both baselines are obtained using their official implementation. IDSplat outperforms the baselines for both settings. **First**, **second**, **third**.

		Anno. free	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	DPSNR $\uparrow$
DeSiRe-GS setting	MARS	$\times$	26.61	-	-	22.21
	SplatAD	$\times$	30.80	0.900	0.160	28.97
	PVG	$\checkmark$	29.77	-	-	27.19
	EmerNeRF	$\checkmark$	25.14	-	-	23.49
	S3Gaussian	$\checkmark$	27.44	-	-	22.92
	DeSiRe-GS	$\checkmark$	28.76	0.873	0.193	26.26
	IDSplat (ours)	$\checkmark$	30.83	0.900	0.160	29.20
	StreetGS	$\times$	33.97	0.926	0.227	28.50
	4DGS	$\times$	34.64	0.940	0.244	29.77
	SplatAD	$\times$	34.24	0.925	0.246	29.68
AD-GS setting	SplatAD (CasTrack)	$\times$	32.52	0.924	0.241	25.31
	PVG	$\checkmark$	29.54	0.895	0.266	21.56
	EmerNeRF	$\checkmark$	31.32	0.881	0.301	21.80
	Grid4D	$\checkmark$	32.19	0.921	0.253	22.77
	AD-GS	$\checkmark$	33.91	0.927	0.228	27.41
	IDSplat (ours)	$\checkmark$	34.59	0.929	0.235	29.63
	CoDa-4DGS	$\checkmark$	28.66	0.900	0.058	-
	IDSplat (ours)	$\checkmark$	30.50	0.875	0.090	-
	SplatFlow	$\checkmark$	28.71	0.874	0.239	-
	IDSplat (ours)	$\checkmark$	29.95	0.879	0.183	-

Table 2. NVS results for lidar point cloud rendering. Our method obtains similar lidar rendering performance to the annotation-based SplatAD.

	Depth $\downarrow$	Intensity $\downarrow$	Drop acc. $\uparrow$	CD $\downarrow$
SplatAD	0.01	0.055	87.3	0.98
IDSplat (ours)	0.01	0.056	87.5	1.24

periments on PandaSet [37]. We further analyze the optimized trajectories using standard tracking metrics. Finally, we perform ablation studies of individual components to understand their impact on overall performance.

4.1. Implementation

IDSplat is implemented in *neurad-studio* [29], built upon the rasterization framework from SplatAD [7, 8]. For pre-processing, we use Grounded-SAM-2 [23] to generate instance masks and DINOv3 [24] for image features. We optimize IDSplat for 30,000 iterations, following the hyperparameter settings in [8] unless otherwise specified. All experiments were run on an NVIDIA A100 40GB GPU. See Sec. A for further details.

4.2. Dataset

Our experiments are conducted on three subsets of Waymo Open Dataset. One is the set of eight sequences used in StreetGS [39] and AD-GS [10], the second is the Waymo

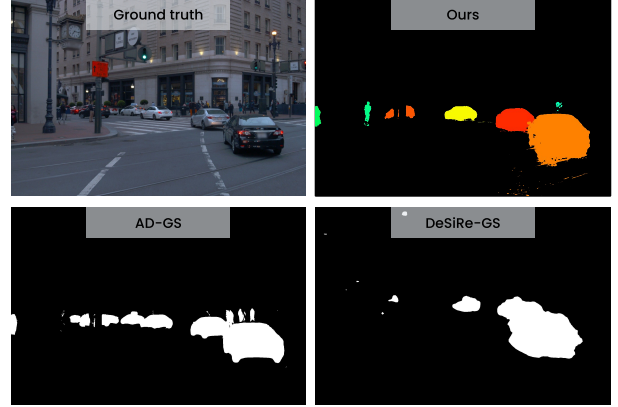


Figure 3. Dynamic mask rendering results. Beyond separating dynamic and static components, our method also renders instance masks for each dynamic object.

Table 3. NVS results under varying view densities. Even with limited training views, our method achieves high DPSNR which highlights the effectiveness of our instance-decomposed model of dynamic objects. **First**, **second**, **third**.

	25%		50%		75%		100%	
	PSNR $\uparrow$	DPSNR $\uparrow$	PSNR $\uparrow$	DPSNR $\uparrow$	PSNR $\uparrow$	DPSNR $\uparrow$	PSNR $\uparrow$	DPSNR $\uparrow$
DeSiRe-GS	24.37	22.97	28.78	27.34	30.04	28.04	35.11	34.99
AD-GS	26.21	22.33	29.97	26.85	30.76	28.07	34.42	35.09
IDSplat (ours)	26.83	26.35	29.19	28.74	30.11	29.25	35.04	33.67

NeRF-On-The-Road (NOTR) dataset, a curated subset of challenging sequences provided in EmerNeRF [40], and the last the set of 4 sequences used in PVG [3]. NOTR contains 32 static, 32 dynamic, and 56 diverse sequences that cover various weather conditions and road types. See Sec. B for more details.

4.3. Baseline comparisons

We compare IDSplat with state-of-the-art self-supervised rendering methods designed for dynamic automotive scenes, focusing our comparison with the following best performing methods, DeSiRe-GS [20], AD-GS [10], CoDa-4DGS [26], and SplatFlow [28]. In addition, we also report results for several supervised methods that rely on bounding box annotations to model dynamic objects. Among these, SplatAD [7] is the most closely related to our approach and serves as a strong reference for assessing how closely IDSplat approaches supervised performance. For further reference, we run SplatAD using tracks from CasTrack [35], a high-performing tracker on the Waymo 3D tracking leaderboard.

NVS results: We report PSNR, SSIM [33], and LPIPS [45] as our primary evaluation metrics, with LPIPS computed using the VGG network [25]. To specifically assess performance in dynamic regions, we also compute dynamic



Figure 4. Qualitative comparisons of novel view synthesis over different view densities (25%, 50%, and 75% of training frames) on the dynamic subset of Waymo NOTR. Our instance-decomposed representation enables high-quality rendering of dynamic objects even when trained with sparse viewpoints.

PSNR (DPSNR) using dynamic object masks. Since both IDSplat and SplatAD support lidar rendering, we additionally evaluate lidar metrics for these methods, reporting median squared depth error, RMSE intensity error, ray drop accuracy, and chamfer distance (CD).

For a fair comparison with prior work, we adopt the data splits and evaluation protocols of each baseline. For DeSiRe-GS, we follow their setup of using 90% of the frames for optimization and evaluating on every tenth frame. We conduct these experiments on the dynamic NOTR subset, using the three front cameras. Due to computational constraints of DeSiRe-GS, we use half-resolution images and only the first 50 frames in each sequence. For AD-GS, we follow their protocol of using 75% of the frames for optimization and evaluating on every fourth frame. Consistent with their settings, we use the eight sequences presented in StreetGS, the front camera, and full-resolution images. We use the same setting as DeSiRe-GS for CoDa-4DGS, but run the experiments on the complete NOTR dataset. To compare our results to SplatFlow, we use the same 4 Waymo sequences used in [3] using 75% of frames for training. To calculate DPSNR for each setting, we use the same dynamic masks originally used by each respective method as the ground truth.

We present comparisons with baseline methods in Tab. 1. IDSplat achieves competitive or superior performance to prior work on both full-frame and dynamic-region evalu-

ations. Notably, IDSplat performs comparably to its supervised counterpart, SplatAD, and achieves similar performance on LiDAR metrics (Tab. 2). IDSplat further outperforms SplatAD when the latter uses tracks from CasTrack, demonstrating that our approach can outperform dataset-specific trackers.

Decomposition quality: To better understand the differences in DPSNR across methods, Fig. 3 visualizes the dynamic object masks generated by DeSiRe-GS, AD-GS, and our method. AD-GS yields reasonably accurate dynamic regions, while DeSiRe-GS often over-segments static areas. Both methods only segment dynamic regions and cannot decompose individual object instances. In contrast, IDSplat generates instance masks that closely align with ground truth. Since our model maintains a consistent set of Gaussians for each object, the appearance and geometry of each actor remain stable, leading to more coherent reconstructions and high rendering quality in dynamic regions. Further, Fig. 5 illustrates how our instance-decomposition enables targeted modification of objects.

A side effect of our dynamic object model is that instance Gaussians may also represent nearby environmental elements that move with the object such as shadows or adjacent appearance effects. This can also be seen in Fig. 3, where the orange vehicle mask slightly extends into the road surface.

4.4. View density comparisons

To assess the robustness of our approach, we conduct additional experiments under varying view densities. Specifically, we evaluate IDSplat alongside DeSiRe-GS and AD-GS, using 25%, 50%, 75%, and 100% of the frames for optimization, while evaluating on the remaining frames, except for 100% where all frames are also used for evaluation. The frames are selected linearly spaced. To enable a consistent comparison across methods, we adopt a unified setup, using the dynamic subset of NOTR with a single camera at full resolution and the first 50 frames of each sequence. We report PSNR and DPSNR in Tab. 3.

Both DeSiRe-GS and AD-GS perform well in the full reconstruction (100%) setting, demonstrating the effectiveness of their time-varying parameterizations in fitting the data. However, evaluations under sparser settings reveal that these parameterizations struggle with larger interpolation gaps, as they are not explicitly constrained to capture the underlying dynamics. DeSiRe-GS degrades significantly on the sparser settings, especially in dynamic regions. AD-GS exhibits more stable performance across different frame densities, but also suffers from a drop in dynamic regions for more sparse settings. In contrast, the performance of IDSplat in dynamic regions is much more consistent with the full-image results across all view densities, and surpasses the baselines with more than 4.8 PSNR in dynamic regions on the most sparse setting.

This is further illustrated in Fig. 4, which shows qualitative comparisons of novel view reconstruction on the dynamic subset of Waymo NOTR for different view densities. As also observed in quantitative results, IDSplat maintains stable reconstruction quality and consistent geometry even when trained with fewer viewpoints.

4.5. Object class comparisons

For further analysis of our method, we present DPSNR for three different classes of road-users; vehicles, pedestrians, and cyclists. These experiments were run using the 50% setting of the setup presented in Sec. 4.4. Masks of the dynamic regions were obtained by projected 3D bounding box annotations, and filtered based on speed and semantic class. As shown in Tab. 4, IDSplat attains the highest accuracy on vehicles, the only class that is modeled as dynamic instances in our approach. IDSplat also shows strong results on pedestrians and cyclists, even though the rigid motion assumption is only partially valid for these classes.

4.6. Generalization

To assess the generalization and robustness of our approach, we further evaluate it on PandaSet [37], without any hyperparameter tuning. Following [7], we use the same 10 sequences and all six available cameras at full resolution, and compute LPIPS using AlexNet [15]. We perform novel

Table 4. NVS results filtered on different dynamic object classes. IDSplat demonstrates strong performance for vehicles and pedestrians, with cyclists being more challenging. First, second, third.

	PSNR	DPSNR _{Vehicle}	DPSNR _{Pedestrian}	DPSNR _{Cyclist}
DeSiRe-GS	28.78	25.04	27.65	29.34
AD-GS	29.97	26.80	27.22	26.52
IDSplat (ours)	29.19	29.02	28.31	27.23

Table 5. NVS results on PandaSet, using all six cameras at full-resolution. IDSplat achieves performance on par with annotation-based methods.

	Anno. free	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$
PVG	$\times$	24.01	0.712	0.452
Street-GS	$\times$	24.73	0.745	0.314
SplatAD	$\times$	26.76	0.815	0.193
IDSplat (ours)	$\checkmark$	26.78	0.814	0.174

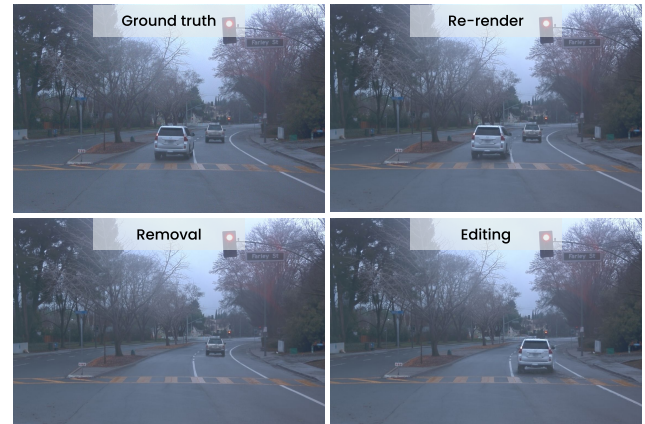


Figure 5. Instance editing. Our instance-decomposed representation enables targeted modifications of individual instances, such as their complete removal or the editing of their trajectories.

view synthesis using 50% of the frames for optimization and the remaining for evaluation, and compare our results against prior state-of-the-art methods. As shown in Tab. 5, IDSplat achieves performance competitive with the best-performing method, SplatAD, despite not relying on any annotations.

4.7. Ablations

We ablate the effectiveness of key components in our method by analyzing their impact on NVS performance, using the AD-GS setting described in Sec. 4.3 but with only 50% of views used for optimization. The results of the ablations are presented in Tab. 6. We observe that the initial filtering of lidar points using eroded masks (a) and DBSCAN

Table 6. NVS results when removing different model components.

		PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	DPSNR $\uparrow$
	Full model	33.59	0.920	0.240	29.35
a)	Eroded masks	33.50	0.920	0.241	29.32
b)	DBSCAN	33.48	0.920	0.242	28.07
c)	Registration	32.60	0.915	0.247	26.67
d)	Image features	32.7	0.917	0.243	25.70
e)	Smoothing	32.96	0.918	0.244	28.01
f)	Outlier rejection	33.38	0.919	0.242	29.44
g)	Refinement	33.32	0.918	0.245	27.36

Table 7. NVS results from using image features from different models when establishing point correspondences. Gray row marks the image features used in our method.

	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	DPSNR $\uparrow$
RGB	32.76	0.917	0.243	25.72
SAM2	33.24	0.918	0.245	27.74
DINOv3 layer 7	33.59	0.920	0.240	29.35
DINOv3 layer 9	33.57	0.919	0.240	29.17
DINOv3 layer 12	33.52	0.920	0.242	28.95

(b) helps improve the final rendering results. Removing the RANSAC-based registration of points (c), and instead naively estimating pose translations by the point clouds centroids, has a severe impact on the performance. Selecting point correspondences using RGB colors instead of image features (d) also degrades the performance notably. Further, we see that the smoothing (e) is a key component for attaining good results in dynamic regions, while the outlier rejection from a first initial smoothing iteration (f) has a small impact on the final rendering results. Finally, we note the importance of refining the trajectories using gradients from the rendering losses during optimization (g). Additionally, we analyze the impact of image features in Tab. 7. Specifically, we run the registration using different layers of DINOv3, image features from SAM2 [22], and RGB colors. We observe that DINOv3 has the most descriptive features for successful registration, with earlier layers giving a slight increase in performance in dynamic regions.

4.8. Tracking

While IDSplat primarily targets high-quality scene reconstruction and rendering, we also evaluate the quality of the resulting motion trajectories. Specifically, we compare our optimized trajectories against human-annotated ground truth trajectories on a combined set of the 32 sequences from the dynamic subset of NOTR and eight sequences from the AD-GS split. We report Multiple Object Tracking Accuracy (MOTA) [1] and Multiple Object Tracking Precision (MOTP) [1] across different distance thresh-

olds, along with detailed metrics such as false positives, ID switches, recall, and precision. The results are provided in Sec. D along with qualitative examples and evaluation details. Since none of the baseline methods perform instance decomposition, tracking evaluation is not applicable to them. Nevertheless, we include our tracking results to establish a reference point and to encourage future research in this direction.

Our trajectories exhibit occasional ID switches, as IDSplat does not explicitly handle identity association across frames. Additional errors arise from stationary objects or incorrectly classified masks generated by Grounded-SAM-2. Furthermore, this evaluation does not account for potential constant offsets between the predicted and ground-truth trajectories, which can occur even under perfect motion tracking, due to partial or incomplete point cloud representations of objects.

4.9. Limitations

While IDSplat provides strong reconstruction performance and enables instance-level decomposition, several limitations remain. First, our object initialization and trajectory estimation are dependent on lidar measurements. Therefore, objects that fall outside of the lidar field of view but are still visible in the cameras are excluded and considered a part of the static scene. Second, our framework assumes that dynamic actors behave as rigid bodies. This works well for vehicles but is less suitable for highly deformable classes such as pedestrians and cyclists, which can lead to reduced reconstruction quality for these categories. Third, because each dynamic object is represented by a single set of Gaussians, nearby environmental effects such as shadows or reflections, which occur consistently in the scene and move with the object, may be represented by the instance. Finally, the current system lacks explicit track management mechanisms, as we do not merge overlapping tracks or handle ID switches, which may result in duplicated or incomplete instances in challenging scenarios. Addressing these limitations presents a promising direction for future work.

5. Conclusion

We presented IDSplat, a zero-shot framework for instance decomposition and neural rendering in dynamic driving scenes. Extensive experiments show that IDSplat achieves competitive or superior performance compared to state-of-the-art self-supervised approaches, and even matches the performance of annotation-based baselines. By representing each actor as a rigid instance, the method establishes a clear separation not only between dynamic and static scene elements, but also between individual dynamic objects, enabling precise scene editing such as object repositioning or removal.

Acknowledgements

We thank Adam Lilja and William Ljungbergh for valuable feedback. This work was partially supported by the Wallenberg AI, Autonomous Systems and Software Program (WASP) funded by the Knut and Alice Wallenberg Foundation. This work was also partially supported by Vinnova, the Swedish Innovation Agency. Computational resources were provided by NAISS at [NSC Berzelius](#), partially funded by the Swedish Research Council, grant agreement no. 2022-06725.

References

- [1] Keni Bernardin and Rainer Stiefelhagen. Evaluating multiple object tracking performance: the CLEAR MOT metrics. *EURASIP Journal on Image and Video Processing*, 2008(1): 246309, 2008. 8, 4, 5
- [2] Yun Chen, Jingkan Wang, Ze Yang, Sivabalan Manivasagam, and Raquel Urtasun. G3R: Gradient guided generalizable reconstruction. In *Computer Vision – ECCV 2024*, page 305–323, Cham, 2025. Springer Nature Switzerland. 2
- [3] Yurui Chen, Chun Gu, Junzhe Jiang, Xiatian Zhu, and Li Zhang. Periodic vibration gaussian: Dynamic urban scene reconstruction and real-time rendering. *International Journal of Computer Vision*, 2026. 1, 2, 5, 6
- [4] Ziyu Chen, Jiawei Yang, Jiahui Huang, Riccardo de Lutio, Janick Martinez Esturo, Boris Ivanovic, Or Litany, Zan Gojic, Sanja Fidler, Marco Pavone, Li Song, and Yue Wang. OmniRe: Omni urban scene reconstruction. In *The Thirteenth International Conference on Learning Representations*, 2025. 1, 2
- [5] Frank Dellaert and GTSAM Contributors. [borglab/gtsam](#), 2022. 4
- [6] Martin Ester, Hans-Peter Kriegel, Jörg Sander, Xiaowei Xu, et al. A density-based algorithm for discovering clusters in large spatial databases with noise. In *Second International Conference on Knowledge Discovery and Data Mining (KDD'96)*, pages 226–231, 1996. 3
- [7] Georg Hess and Carl Lindström. [splatad](#). <https://github.com/carlinds/splatad>, 2025. 5, 7, 1
- [8] Georg Hess, Carl Lindström, Maryam Fatemi, Christoffer Petersson, and Lennart Svensson. SplatAD: Real-time lidar and camera rendering with 3D Gaussian splatting for autonomous driving. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 11982–11992, 2025. 1, 2, 3, 4, 5
- [9] Nan Huang, Xiaobao Wei, Wenzhao Zheng, Pengju An, Ming Lu, Wei Zhan, Masayoshi Tomizuka, Kurt Keutzer, and Shanghang Zhang. S³gaussian: Self-supervised street gaussians for autonomous driving. *arXiv preprint arXiv:2405.20323*, 2024. 1, 2
- [10] Xu Jiawei, Deng Kai, Fan Zexin, Wang Shenlong, Xie Jin, and Yang Jian. AD-GS: Object-aware B-Spline Gaussian splatting for self-supervised autonomous driving. *International Conference on Computer Vision*, 2025. 1, 2, 5
- [11] Bernhard Kerbl, Georgios Kopanas, Thomas Leimkühler, and George Drettakis. 3D Gaussian splatting for real-time radiance field rendering. *ACM Transactions on Graphics*, 42(4), 2023. 1, 2
- [12] Mustafa Khan, Hamidreza Fazlali, Dhruv Sharma, Tongtong Cao, Dongfeng Bai, Yuan Ren, and Bingbing Liu. AutoSplat: Constrained gaussian splatting for autonomous driving scene reconstruction. *arXiv preprint arXiv:2407.02598*, 2024. 2
- [13] Shakiba Kheradmand, Daniel Rebain, Gopal Sharma, Weiwei Sun, Yang-Che Tseng, Hossam Isack, Abhishek Kar, Andrea Tagliasacchi, and Kwang Moo Yi. 3D Gaussian splatting as markov chain monte carlo. In *Advances in Neural Information Processing Systems*, 2024. 4, 1, 2
- [14] Diederik P Kingma. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014. 2
- [15] Alex Krizhevsky, Ilya Sutskever, and Geoffrey E Hinton. ImageNet classification with deep convolutional neural networks. *Advances in Neural Information Processing Systems*, 25, 2012. 7
- [16] Carl Lindström, Georg Hess, Adam Lilja, Maryam Fatemi, Lars Hammarstrand, Christoffer Petersson, and Lennart Svensson. Are NeRFs ready for autonomous driving? Towards closing the real-to-simulation gap. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4461–4471, 2024. 1
- [17] William Ljungbergh, Adam Tonderski, Joakim Johnander, Holger Caesar, Kalle Åström, Michael Felsberg, and Christoffer Petersson. NeuroNCAP: Photorealistic closed-loop safety testing for autonomous driving. In *European Conference on Computer Vision*, pages 161–177. Springer, 2024. 1
- [18] Yunxuan Mao, Rong Xiong, Yue Wang, and Yiyi Liao. Unire: Unsupervised instance decomposition for dynamic urban scene reconstruction, 2025. 2
- [19] Ben Mildenhall, Pratul P Srinivasan, Matthew Tancik, Jonathan T Barron, Ravi Ramamoorthi, and Ren Ng. NeRF: Representing scenes as neural radiance fields for view synthesis. *Communications of the ACM*, 65(1):99–106, 2021. 2
- [20] Chensheng Peng, Chengwei Zhang, Yixiao Wang, Chenfeng Xu, Yichen Xie, Wenzhao Zheng, Kurt Keutzer, Masayoshi Tomizuka, and Wei Zhan. DeSiRe-GS: 4D street gaussians for static-dynamic decomposition and surface reconstruction for urban driving scenes. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 6782–6791, 2025. 1, 2, 5, 4
- [21] Mahan Rafidashti, Ji Lan, Maryam Fatemi, Junsheng Fu, Lars Hammarstrand, and Lennart Svensson. NeuRadar: Neural radiance fields for automotive radar point clouds. In *2025 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, pages 2479–2489, 2025. 2
- [22] Nikhila Ravi, Valentin Gabeur, Yuan-Ting Hu, Ronghang Hu, Chaitanya Ryali, Tengyu Ma, Haitham Khedr, Roman Rädle, Chloe Rolland, Laura Gustafson, Eric Mintun, Junting Pan, Kalyan Vasudev Alwala, Nicolas Carion, Chao-Yuan Wu, Ross Girshick, Piotr Dollár, and Christoph Feichtenhofer. SAM 2: Segment anything in images and videos. *arXiv preprint arXiv:2408.00714*, 2024. 8

- [23] Tianhe Ren and Shuo Shen. Grounded-SAM-2. <https://github.com/IDEA-Research/Grounded-SAM-2>, 2024. 2, 3, 5
- [24] Oriane Siméoni, Huy V. Vo, Maximilian Seitzer, Federico Baldassarre, Maxime Oquab, Cijo Jose, Vasil Khalidov, Marc Szafraniec, Seungeun Yi, Michaël Ramamonjisoa, Francisco Massa, Daniel Haziza, Luca Wehrstedt, Jianyuan Wang, Timothée Darcet, Théo Moutakanni, Leonel Sentana, Claire Roberts, Andrea Vedaldi, Jamie Tolan, John Brandt, Camille Couprie, Julien Mairal, Hervé Jégou, Patrick Labatut, and Piotr Bojanowski. DINOv3. *arXiv preprint arXiv:2508.10104*, 2025. 2, 3, 5
- [25] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014. 5
- [26] Rui Song, Chenwei Liang, Yan Xia, Walter Zimmer, Hu Cao, Holger Caesar, Andreas Festag, and Alois Knoll. Coda-4dgs: Dynamic gaussian splatting with context and deformation awareness for autonomous driving. In *IEEE/CVF International Conference on Computer Vision (ICCV)*. IEEE/CVF, 2025. 2, 5
- [27] Pei Sun, Henrik Kretschmar, 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, 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)*, 2020. 2, 4
- [28] Su Sun, Cheng Zhao, Zhuoyang Sun, Yingjie Victor Chen, and Mei Chen. SplatFlow: Self-supervised dynamic gaussian splatting in neural motion flow field for autonomous driving. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 27487–27496, 2025. 2, 5
- [29] Adam Tonderski, Carl Lindström, Georg Hess, and William Ljungbergh. neurad-studio. <https://github.com/georghess/neurad-studio>, 2024. 5
- [30] Adam Tonderski, Carl Lindström, Georg Hess, William Ljungbergh, Lennart Svensson, and Christoffer Petersson. NeuRAD: Neural rendering for autonomous driving. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14895–14904, 2024. 1, 2
- [31] Haithem Turki, Jason Y Zhang, Francesco Ferroni, and Deva Ramanan. SUDS: Scalable urban dynamic scenes. In *Computer Vision and Pattern Recognition (CVPR)*, 2023. 1, 2
- [32] Shinji Umeyama. Least-squares estimation of transformation parameters between two point patterns. *IEEE Transactions on pattern analysis and machine intelligence*, 13(4): 376–380, 2002. 3
- [33] Zhou Wang, A.C. Bovik, H.R. Sheikh, and E.P. Simoncelli. Image quality assessment: from error visibility to structural similarity. *IEEE Transactions on Image Processing*, 13(4): 600–612, 2004. 5
- [34] Guanjun Wu, Taoran Yi, Jiemin Fang, Lingxi Xie, Xiaopeng Zhang, Wei Wei, Wenyu Liu, Qi Tian, and Xinggang Wang. 4D Gaussian splatting for real-time dynamic scene rendering. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 20310–20320, 2024. 2
- [35] Hai Wu, Wenkai Han, Chenglu Wen, Xin Li, and Cheng Wang. 3d multi-object tracking in point clouds based on prediction confidence-guided data association. *IEEE Transactions on Intelligent Transportation Systems*, 23(6):5668–5677, 2022. 5
- [36] Zirui Wu, Tianyu Liu, Liyi Luo, Zhide Zhong, Jianteng Chen, Hongmin Xiao, Chao Hou, Haozhe Lou, Yuan-tao Chen, Runyi Yang, Yuxin Huang, Xiaoyu Ye, Zike Yan, Yongliang Shi, Yiyi Liao, and Hao Zhao. MARS: An instance-aware, modular and realistic simulator for autonomous driving. In *Artificial Intelligence*, pages 3–15, Singapore, 2024. Springer Nature Singapore. 1, 2
- [37] Pengchuan Xiao, Zhenlei Shao, Steven Hao, Zishuo Zhang, Xiaolin Chai, Judy Jiao, Zesong Li, Jian Wu, Kai Sun, Kun Jiang, Yunlong Wang, and Diange Yang. PandaSet: Advanced sensor suite dataset for autonomous driving. In *2021 IEEE International Intelligent Transportation Systems Conference (ITSC)*, pages 3095–3101, 2021. 5, 7
- [38] Jiawei Xu, Zexin Fan, Jian Yang, and Jin Xie. Grid4D: 4d decomposed hash encoding for high-fidelity dynamic gaussian splatting. *Advances in Neural Information Processing Systems*, 37:123787–123811, 2024. 2
- [39] Yunzhi Yan, Haotong Lin, Chenxu Zhou, Weijie Wang, Haiyang Sun, Kun Zhan, Xianpeng Lang, Xiaowei Zhou, and Sida Peng. Street gaussians for modeling dynamic urban scenes. In *Computer Vision – ECCV 2024*, 2024. 1, 2, 5
- [40] Jiawei Yang, Boris Ivanovic, Or Litany, Xinshuo Weng, Seung Wook Kim, Boyi Li, Tong Che, Danfei Xu, Sanja Fidler, Marco Pavone, et al. EmerNeRF: Emergent spatial-temporal scene decomposition via self-supervision. *arXiv preprint arXiv:2311.02077*, 2023. 1, 2, 5
- [41] Jiawei Yang, Jiahui Huang, Yuxiao Chen, Yan Wang, Boyi Li, Yurong You, Maximilian Igl, Apoorva Sharma, Peter Karkus, Danfei Xu, Boris Ivanovic, Yue Wang, and Marco Pavone. STORM: Spatio-temporal reconstruction model for large-scale outdoor scenes. *arXiv preprint arXiv:2501.00602*, 2025. 2
- [42] Ze Yang, Yun Chen, Jingkan Wang, Sivabalan Manivasagam, Wei-Chiu Ma, Anqi Joyce Yang, and Raquel Urtasun. UniSim: A neural closed-loop sensor simulator. In *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1389–1399, 2023. 1, 2
- [43] Ziyi Yang, Xinyu Gao, Wen Zhou, Shaohui Jiao, Yuqing Zhang, and Xiaogang Jin. Deformable 3D Gaussians for high-fidelity monocular dynamic scene reconstruction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 20331–20341, 2024. 2
- [44] Yuanwen Yue, Anurag Das, Francis Engelmann, Siyu Tang, and Jan Eric Lenssen. Improving 2D feature representations by 3D-aware fine-tuning. In *Computer Vision – ECCV 2024*, page 57–74, Cham, 2025. Springer Nature Switzerland. 2

- [45] Richard Zhang, Phillip Isola, Alexei A. Efros, Eli Shechtman, and Oliver Wang. The unreasonable effectiveness of deep features as a perceptual metric. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018. [5](#)
- [46] Ruida Zhang, Chengxi Li, Chenyangguang Zhang, Xingyu Liu, Haili Yuan, Yanyan Li, Xiangyang Ji, and Gim Hee Lee. Street gaussians without 3D object tracker. *arXiv preprint arXiv:2412.05548*, 2024. [1](#), [2](#)
- [47] Hongyu Zhou, Jiahao Shao, Lu Xu, Dongfeng Bai, Weichao Qiu, Bingbing Liu, Yue Wang, Andreas Geiger, and Yiyi Liao. HUGS: Holistic urban 3d scene understanding via gaussian splatting. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 21336–21345, 2024. [2](#)
- [48] Xiaoyu Zhou, Zhiwei Lin, Xiaojun Shan, Yongtao Wang, Deqing Sun, and Ming-Hsuan Yang. DrivingGaussian: Composite gaussian splatting for surrounding dynamic autonomous driving scenes. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 21634–21643, 2024. [1](#), [2](#)