

RadarGen: Automotive Radar Point Cloud Generation from Cameras

Tomer Borreda¹, Fangqiang Ding^{1,2}, Sanja Fidler^{3,4,5},
Shengyu Huang³, and Or Litany^{1,3}

¹ Technion, Haifa, Israel

² MIT, Cambridge MA, USA

³ NVIDIA, Santa Clara CA, USA

⁴ University of Toronto, Toronto, Canada

⁵ Vector Institute, Toronto, Canada

<https://radargen.github.io/>

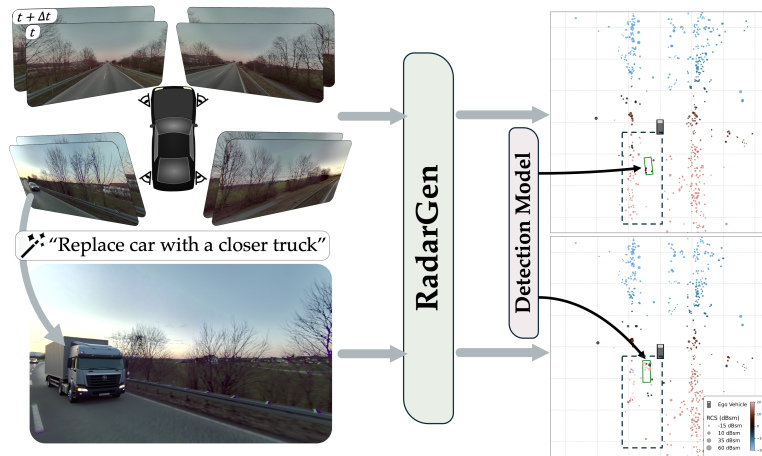


Fig. 1: Controllable radar synthesis from vision. (Top) Given multi-view camera images, *RadarGen* generates realistic radar point clouds that align with real-world radar statistics and can be consumed by downstream perception models. (Bottom) The generation is semantically consistent: modifying the input scene with an off-the-shelf image editing tool (e.g., replacing a distant car with a closer truck) updates the radar response, removing returns from newly occluded regions and reflecting the new object geometry.

Abstract. We present RadarGen, a diffusion model for synthesizing realistic automotive radar point clouds from multi-view camera imagery. RadarGen adapts efficient image-latent diffusion to the radar domain by representing radar measurements in bird’s-eye-view form that encodes spatial structure together with radar cross section (RCS) and Doppler attributes. A lightweight recovery step reconstructs point clouds from the generated maps. To better align generation with the visual scene, RadarGen incorporates BEV-aligned depth, semantic, and motion cues extracted from pretrained foundation models, which guide the stochastic

generation process toward physically plausible radar patterns. Conditioning on images makes the approach broadly compatible, in principle, with existing visual datasets and simulation frameworks, offering a scalable direction for multimodal generative simulation. Evaluations on large-scale driving data show that RadarGen captures characteristic radar measurement distributions and reduces the gap to perception models trained on real data, marking a step toward unified generative simulation across sensing modalities.

1 Introduction

Recent advances in neural and generative simulation have made it increasingly practical to synthesize photorealistic data at scale for autonomous driving. By reconstructing real scenes with neural fields or generating entirely new ones using video diffusion models, these systems can produce diverse and controllable environments that closely mimic real sensor observations. This capability enables large scale resimulation of traffic, lighting, and weather conditions without costly rerecording or manual setup [42, 50, 61]. Despite this rapid progress, most neural simulators remain limited to the visual domain, focusing on the generation of RGB imagery and video. Recent efforts have begun extending these ideas to LiDAR, demonstrating controllable three-dimensional point cloud generation from camera inputs [56, 59, 89]. Radar, however, remains an open frontier. Although it is already ubiquitous in production vehicles, providing low cost, lightweight, and weather resilient perception, it has received far less attention within the generative modeling community. This imbalance limits the fidelity of current neural simulators, which cannot reproduce radar’s distinctive sensing characteristics, including signal sparsity, radar cross section (RCS), and Doppler.

Generating radar data poses unique challenges. Radar measurements exhibit strong stochasticity due to multipath reflections, interference, and material-dependent scattering that vary with scene geometry. Operating at longer wavelengths, radar interacts with surfaces and internal structures beyond what cameras or LiDAR perceive, making it highly complementary yet difficult to model from vision alone. A further challenge lies in the nature of available radar data. In most large scale driving datasets, radar is provided only after proprietary signal processing that converts raw radio frequency waveforms into sparse point clouds with RCS and Doppler values. This processing chain, which includes range-Doppler transforms, beamforming, and detection algorithms such as constant false alarm rate (CFAR), is closed and lossy, discarding phase and other fine grained signal information. In practice, storing raw radar signals is extremely memory intensive, so even commercial survey vehicles often record only processed point clouds. As a result, we opted for a point cloud representation as it remains the practical representation of radar data, providing diversity and scalability critical for generative model training.

To address these challenges, we propose *RadarGen*, a generative framework for synthesizing automotive radar point clouds directly from camera imagery.

RadarGen learns a distribution over radar observations conditioned on the visual scene, producing diverse and semantically consistent measurements rather than a single deterministic prediction, reflecting the inherent stochasticity of real radar signals. Conditioning on images allows *RadarGen* to leverage existing visual data and simulators, providing a scalable and modular way to enrich them with realistic radar so that radar-dependent perception systems can operate on camera-only and edited scenes, without re-recording radar.

A key design choice in *RadarGen* is to build on SANA [82], an efficient image-latent diffusion model proven effective and scalable for image synthesis. To the best of our knowledge, no existing generative model can reliably support scene level point cloud synthesis from multi-view images of real driving scenes, making an image diffusion backbone a natural and practical foundation. SANA’s architecture supports conditioning on visual input while efficiently handling the large number of tokens introduced by the multi-channel radar representation and by the additional conditioning cues incorporated later in our framework, allowing *RadarGen* to remain computationally efficient and deployable in large scale simulation settings.

To make radar data compatible with the image diffusion backbone, we express each radar point cloud as an image-like bird’s eye view (BEV) representation. This representation encodes spatial structure together with radar cross section (RCS) and Doppler attributes, enabling all radar channels to share a unified latent space and allowing us to reuse SANA’s pretrained autoencoder without modification. A lightweight smoothing and recovery step preserves geometric accuracy while ensuring stable latent encoding.

Learning radar purely from images is inherently difficult, as it requires reasoning about depth, semantics, and motion. To avoid forcing the denoiser to learn these cues from scratch, *RadarGen* leverages pretrained foundation models that provide dense visual priors. Depth, semantic, and motion cues extracted from each camera view are projected into bird’s eye view to align with the radar representation and are fused within the diffusion model.

Since no established benchmark exists for radar point cloud generation, we introduce a dedicated suite of metrics that measures geometric fidelity, radar attribute fidelity, and distribution similarity, both for background and foreground points. We hope this benchmark will support future progress in radar generation.

We summarize our main contributions below.

- We present *RadarGen*, the first probabilistic diffusion framework to generate realistic automotive radar point clouds including location, RCS, and Doppler from multi-view camera inputs.
- We introduce a latent diffusion methodology that trains on a BEV image representation of sparse radar attributes, conditioned by BEV-aligned visual depth, semantic, and motion priors from foundation models.
- We establish comprehensive metrics for radar point clouds based on geometric fidelity, radar attribute fidelity, and distribution similarity, validating our design through extensive ablations.

- We demonstrate that the generated radar data can be interpreted by detectors trained on real data and supports applications such as scene editing.

2 Related Work

2.1 Physics-based radar simulation

Physics-based simulators model the emission, propagation and reception of electromagnetic (EM) waves based on physical laws. To rigorously simulate time-domain EM propagation, some works directly solve Maxwell’s equations in the integral [10] or differential [25, 43, 70] form, offering high physical accuracy but are computation-intensive for real-time applications. For practical usage, other works approximate EM wave propagation based on geometric optics and ray-tracing [85]. This technique has been widely adapted to enhance system-level fidelity [22, 28, 32, 73], improve scalability and real-time performance [30, 67, 71], and support specific modern applications [3, 27, 66]. Commercial software also implement similar techniques [21, 58]. RadSimReal [5] utilized *standard* 3D reflection signals to lessen dependence on radar-specific internals in ray-tracing simulators. Distinct from ray-tracing, graphics-based simulators [51, 65, 69] exploit the rasterization pipeline of graphics engines as an efficient physics proxy, generating radar measurements from depth in real time. However, the aforementioned physical simulators rely on manually created assets, demanding substantial engineering effort to cover long-tail conditions.

2.2 Data-driven radar simulation

The heavy engineering efforts of physics-based radar simulators have motivated data-driven alternatives. One line adapts NeRF [6, 34, 40, 46, 55] or 3D Gaussian Splatting [36, 38, 41] to radar, learning scene-specific representations for novel-view rendering. These reconstructions, however, require multi-view radar captures per scene and transfer poorly to unseen scenes. In contrast, generative radar simulation methods synthesize measurements directly from conditioning inputs, enabling controllable generation for novel, unobserved environments. Prior works employ GANs [26] and VAEs [37] conditioned on object distance [24], scene layouts [16, 77], elevation maps [76], or LiDAR [41], but they ignore images as conditioning.

Two recent works [11, 17] use visual conditioning for radar generation, but focus on human-centric scenarios and depend on explicit physical modeling of signal-scene interactions, which limits scalability. For autonomous driving, Xiao et al. [81] synthesize radar cubes with a U-Net conditioned on camera/LiDAR plus waveform-parameter embeddings, while Rangaraj et al. [57] generate range-azimuth maps using an autoencoder conditioned on depth and segmentation. Yet these methods target radar raw data that is unavailable for large-scale training [20]. Closer to our setting, Song et al. [68] and Alkanat et al. [2] produce radar point clouds from LiDAR/RGB with convolutional networks, but their deterministic mappings under-model radar stochasticity and also do not exploit large pretrained foundation models for comprehensive scene encoding.

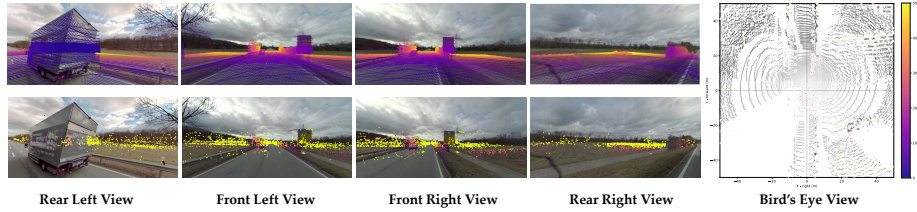


Fig. 2: Comparison of LiDAR and radar point clouds in camera and bird’s eye views. In the camera view, LiDAR is easier to interpret, whereas radar is more difficult due to its sparsity and lack of structure. In the bird’s eye view, radar points are easier to differentiate and interpret.

2.3 Generative point cloud models

Generative models specific to radar point clouds remain scarce in literature, so we review general point cloud generative frameworks as references. Many works designed unconditional models [1, 9, 44, 47, 75, 84, 86, 88] for point cloud completion and generation or text-conditioned [49, 80] generative frameworks, which rely on large priors and lack spatial grounding and controllability. In contrast, some recent works [39, 45, 74] learn to generate point clouds conditioned on RGB images. However, these works are object-centric, only generating point clouds for object shapes. Instead, some methods are proposed to tackle scene-level generation of LiDAR point clouds with generative techniques such as GAN [7, 62], VQ-VAE [83] and diffusion models [33, 48, 56, 79, 89]. While effective, these methods usually lack image-based conditional design and work on dense LiDAR range images, which cannot transfer to sparse and non-uniformly sampled radar point clouds [18, 19, 52].

3 Preliminaries

We start by reviewing the challenges in generating radar point clouds (Sec. 3.1), and how we can benefit from efficient diffusion models (Sec. 3.2). This sets the stage for our proposed radar point cloud generation model (Sec. 4).

3.1 Challenges in radar point cloud generation

The key challenge in generative modeling of radar point clouds lies in their unique data characteristics. Radar point clouds are sparse, unordered 3D point sets with highly non-uniform sampling. Unlike LiDAR, whose returns are uniform along angular dimensions and can be reshaped into dense range images, radar detections arise from peak-based target extraction (e.g., CFAR [64]), and cannot form dense, grid-aligned measurements. See Fig. 2 and the Supplementary Material for more information. Further, each radar point carries sensor-specific attributes beyond 3D position—Radar Cross-Section (RCS), a proxy for reflectivity, mainly

affected by the object material, geometry and incident angles, and Doppler velocity, the measurement of relative radial velocity. As a result of this sparse, non-grid structure, we cannot effectively represent radar data in a range-image format. We thus adopt a multi-image bird’s-eye-view (BEV) representation as a scalable alternative that is well-suited to the data’s sparse nature.

3.2 Efficient diffusion models

A significant challenge in generative modeling is the synthesis of large-scale, explicit 3D point clouds conditioned on images. As discussed in Sec. 3.1 we use a multi-image BEV representation, which is inherently high-dimensional, posing a significant challenge for standard diffusion architectures. To tackle this high-dimensional generation problem at a scene-level scale, we must leverage an efficient diffusion architecture. Specifically, Latent Diffusion Models (LDMs) [31, 60] which compress images into smaller representations. Standard LDMs are insufficient, as they typically use an autoencoder (AE) with an 8x downsampling factor and a diffusion backbone with $O(N^2)$ quadratic self-attention complexity. We therefore build upon SANA [82], a framework that achieves efficiency by employing an AE with 32x compression and replacing the costly self-attention with an $O(N)$ linear attention mechanism.

4 Method

Problem statement. We address the task of surround-camera to radar point cloud generation. Given scene imagery captured by a rig of N outward-facing cameras $\mathbf{I}^t = \{I_1^t, \dots, I_N^t\}$ at two consecutive timesteps $(t, t + \Delta t)$, with Δt time elapsed between them, together with known camera intrinsics and extrinsics, the goal is to generate a radar point cloud $\mathcal{P}^t = \{(x_i, y_i, r_i, d_i)\}_{i=1}^L$ representing the radar detections at time t in the ego-vehicle coordinate frame. Each element describes the planar spatial coordinates (x, y) , radar cross section r (RCS), and Doppler velocity d .

Method overview. *RadarGen* consists of three main components (See Fig. 3 and 4). During training: (1) the conversion of the *sparse* radar point cloud $\mathcal{P}^t$ into a set of *dense* BEV representations suitable for a latent image diffusion model (Sec. 4.1); and (2) conditioning the generator on geometric and semantic cues derived from pretrained vision models to enhance structural and semantic understanding (Sec. 4.2). With these we learn the conditional distribution $p_\theta(\mathcal{P}^t | \mathbf{I}^t, \mathbf{I}^{t+\Delta t})$, which models the stochastic mapping from multi-view imagery to the radar point cloud at time t . At inference: (3) the recovery of the final sparse radar point cloud from the generated dense BEV predictions (Sec. 4.3).

4.1 Representing radar as images

Our goal is to train a diffusion denoiser that operates in the latent space of a pretrained autoencoder. To keep the autoencoder intact, the radar observations

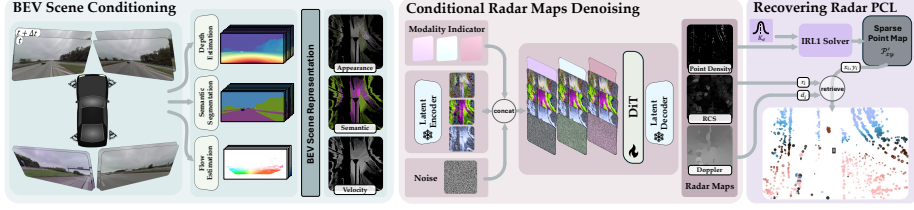


Fig. 3: Overview of RadarGen. (Left) Multi-view posed images at time t and $t + \Delta t$ are fed through foundation models of metric depth estimation [54], semantic segmentation [13], and optical flow [87], enabling projection of the scene to BEV, encoding different information through color. (Middle) Encoded BEV representation is concatenated with a modality indicator specifying which map type to generate. During inference, the map is initialized as noise; during training, noise is added to the GT maps, and the Latent Encoder/Decoder are frozen while SANA’s DiT [53, 82] is fine-tuned. (Right) During inference, the generated Point Density Map is deconvolved using an IRL1 Solver. The resulting sparse map is used to retrieve the RCS and Doppler values at corresponding locations to yield the final generated radar point cloud. Point color represents Doppler and point size represents RCS.

must therefore be represented as image-like signals that fit naturally within its latent distribution. However, radar measurements are sparse point clouds, far from the dense appearance statistics of natural images. We address this gap by transforming each radar point cloud into a set of dense, two-dimensional BEV maps that preserve radar-specific attributes while remaining compatible with the image-based AE. Fig. 4 illustrates the following steps.

Each radar point cloud is first projected onto the BEV plane by discarding elevation, which carries little discriminative information due to radar’s limited vertical resolution in automotive setups. The projected points are then rasterized to form a BEV point map in which detections correspond to occupied pixels.

We construct three single-channel BEV maps: a Point Density Map (M_p), an RCS Map (M_r), and a Doppler Map (M_d). The Point Density Map M_p is obtained by convolving the sparse BEV point map $\mathcal{P}_{xy}$ with a Gaussian kernel K_σ of variance σ , i.e., $M_p = K_\sigma * \mathcal{P}_{xy}$, producing a smooth estimate of radar return density. Increasing σ yields smoother maps with higher reconstruction fidelity under an AE, but also complicates point cloud recovery (see Sec. 4.3). For the RCS and Doppler maps, M_r and M_d , each pixel inherits the attribute value (RCS or Doppler) of the nearest radar detection, yielding a piecewise-constant map defined by the Voronoi tessellation of the detections.

To ensure compatibility with the autoencoder’s RGB input space, each BEV map is replicated across three channels. We then encode each of these images independently, obtaining the latent representations z_p , z_r , and z_d that serve as the supervision targets for training the diffusion denoiser.

Additional preprocessing details, including region-of-interest cropping and grid resolution, are provided in Sec. 4.4.

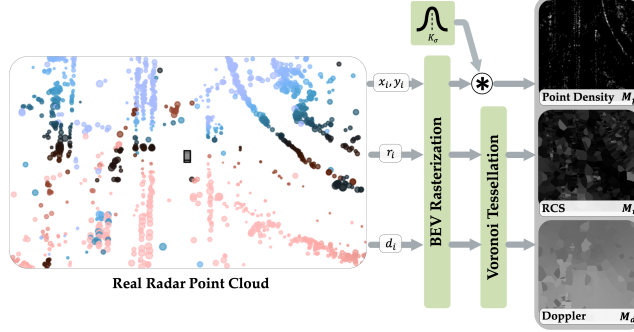


Fig. 4: Overview of representing radar as images (Sec. 4.1). Constructing radar maps from a radar point cloud requires first rasterizing each point to BEV. The point locations are then convolved with a Gaussian kernel K_σ to produce the Point Density Map M_p . A Voronoi tessellation is also constructed, where each cell inherits the RCS and Doppler attributes from its corresponding point, producing the maps M_r and M_d respectively. Point color represents Doppler and point size represents RCS.

4.2 Conditional radar generation

To simplify the conditional radar generation task, we employ pretrained vision foundation models to extract relevant factors from conditional inputs and project them into BEV representations which are semantically rich and spatially aligned with the radar BEV targets.

Specifically, we use a depth estimation network [54] to provide geometric cues, a semantic segmentation network [13] to supply categorical context, and an optical flow model [87] to infer per-pixel motion. These cues are projected and fused into a unified bird’s-eye-view scene representation that serves as the conditioning input to the radar generator.

BEV scene conditioning. The conditioning representation, denoted $\mathbf{c}$, comprises three BEV maps: an Appearance map, a Semantic map, and a Radial Velocity map (see Fig. 3). We first project the N outward-facing camera images $\mathbf{I}^t = \{I_1^t, \dots, I_N^t\}$ into BEV using predicted metric depth [54]. For each image I_k^t , a depth estimation model provides a dense point map $P_{I_k^t}$, which is then transformed into the ego-vehicle coordinate frame using the known camera intrinsics and extrinsics. All point maps are merged and rasterized onto a common BEV grid to form the geometric backbone of the condition.

The points that construct the *Appearance* and *Semantic maps* obtain their color from the original images and segmented images, respectively. Using color-coded semantics, rather than one-hot encodings, ensures that the conditioning maps retain image-like statistics compatible with pretrained image encoders used later in our pipeline.

The *Radial Velocity map* provides motion cues analogous to Doppler velocity. To construct it, we first estimate optical flow between consecutive frames $(I^t, I^{t+\Delta t})$ using a pretrained flow network [87]. Let $\mathbf{f}(x)$ denote the flow vector

at pixel x in I^t . Using the predicted depth at time t and $t + \Delta t$, we backproject x in I^t and its corresponding pixel $x + \mathbf{f}(x)$ in $I^{t+\Delta t}$ to 3D points $p^t(x)$ and $p^{t+\Delta t}(x)$ in the ego-vehicle frame. We then approximate the 3D velocity as $v(x) \approx \frac{p^{t+\Delta t}(x) - p^t(x)}{\Delta t}$, and retain only its component along the radial direction from the ego-vehicle to obtain a Doppler-like value. These radial velocities are rasterized and aggregated into a BEV grid, yielding the Radial Velocity conditioning map.

This BEV scene representation provides pixel-aligned conditioning for the radar diffusion model, ensuring that geometric, semantic, and dynamic cues are spatially consistent with the generated radar BEV maps.

Conditional radar maps denoising. With the radar maps represented in latent space, the learning objective becomes to model the distribution of plausible radar latents conditioned on the visual scene. We therefore train a conditional diffusion denoiser that learns $p(z_p, z_r, z_d \mid \mathbf{c})$, where $\mathbf{c}$ denotes the BEV conditioning maps. Because both the conditioning and target latents are defined in a spatially aligned BEV coordinate frame, the conditioning tensors can be concatenated directly along the channel dimension [35], providing the model with a compact yet expressive geometric prior. This formulation allows the network to generate diverse yet physically consistent radar realizations aligned with the observed camera views.

The denoiser ϵ_θ is implemented as a Diffusion Transformer (DiT) [53] that operates jointly on the latent representations of the three radar maps—density, RCS, and Doppler. During each denoising step, the latents for these maps are concatenated into a single token sequence and processed through shared self-attention, enabling the network to capture correlations across radar modalities. Modality-specific identifiers, defined as learnable embeddings m_i [29] concatenated feature-wise to each radar map $i \in \{p, r, d\}$, guide the transformer to preserve the distinct statistics of each channel while still sharing information between them. At inference, Gaussian noise is iteratively denoised under BEV conditioning to produce the latent radar maps $\{z'_p, z'_r, z'_d\}$, which are then decoded into the final BEV radar images by the pretrained decoder $\mathcal{D}$.

4.3 Recovering the sparse radar point cloud

As described in Sec. 4.1, the point density map is modeled as the convolution of a sparse point map with a Gaussian kernel, $M_p = K_\sigma * \mathcal{P}_{xy}$. The generative model therefore produces a blurred density estimate M'_p , from which the underlying discrete point locations must be recovered.

We formulate recovery as an explicit deconvolution problem, taking advantage of the fact that the blurring kernel K_σ is known and fixed, since it was used to generate the training targets. This knowledge of the forward blurring process enables us to pose a well-defined inverse problem for recovery, yielding a principled and controllable reconstruction of the sparse radar detections.

Table 1: Quantitative evaluation. *RadarGen* broadly outperforms the baseline on geometric fidelity (CD, IoU, Density Similarity, Hit Rate), radar attribute fidelity (DA Recall, Precision, F1), and distribution similarity (MMD). Shown is the mean $\pm$ s.d., computed over all timestamps for the Entire Area and over all foreground objects for the Foreground.

Method	Entire Area										
	CD Loc. ($\downarrow$)	CD Full ($\downarrow$)	IoU@1m ($\uparrow$)	DA Recall ($\uparrow$)	DA Prec. ($\uparrow$)	DA F1 ($\uparrow$)	MMD Loc. ($\downarrow$)	MMD RCS ($\downarrow$)	MMD Doppler ($\downarrow$)		
Baseline	1.84 $\pm$ 0.48	0.038 $\pm$ 0.009	0.23 $\pm$ 0.10	0.15 $\pm$ 0.10	0.14 $\pm$ 0.10	0.14 $\pm$ 0.09	0.368 $\pm$ 0.151	0.36 $\pm$ 0.25	0.65 $\pm$ 0.64		
RadarGen	1.68 $\pm$ 0.39	0.040 $\pm$ 0.008	0.31 $\pm$ 0.11	0.23 $\pm$ 0.12	0.26 $\pm$ 0.12	0.24 $\pm$ 0.12	0.056 $\pm$ 0.062	0.09 $\pm$ 0.15	0.31 $\pm$ 0.74		
Foreground											
	CD Loc. ($\downarrow$)	CD Full ($\downarrow$)	Density Sim. ($\uparrow$)	Hit Rate ($\uparrow$)	MMD Car ($\downarrow$)			MMD Truck ($\downarrow$)		MMD Trailer ($\downarrow$)	
					Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler		
Baseline	1.32 $\pm$ 0.79	0.075 $\pm$ 0.049	0.35 $\pm$ 0.43	0.37	0.035	0.753	0.549	0.167	0.202	0.485	0.0459 0.064 0.607
RadarGen	0.95 $\pm$ 0.65	0.069 $\pm$ 0.049	0.51 $\pm$ 0.41	0.66	0.037	0.006	0.014	0.024	0.031	0.060	0.0069 0.022 0.046

We solve for the sparse point map $\mathcal{P}'_{xy}$ via an L1-regularized, non-negative deconvolution (LASSO) [72]:

$$\min_{\mathcal{P}_{xy} \geq 0} \frac{1}{2} \|K_{\sigma} * \mathcal{P}_{xy} - M'_p\|_2^2 + \lambda \|\mathcal{P}_{xy}\|_1, \quad (1)$$

where λ balances data fidelity and sparsity. We optimize this objective using an Iteratively Reweighted L1 (IRL1) [15] scheme with a FISTA [4] solver, which provides fast convergence and stable recovery of sparse signals. The resulting $\mathcal{P}'_{xy}$ is thresholded to extract the final set of L 2D coordinates $\{(x_i, y_i)\}_{i=1}^L$. For each recovered location, we retrieve the RCS and Doppler attributes from their corresponding map positions, producing the final reconstructed radar point cloud $\mathcal{P}' = \{(x_i, y_i, r_i, d_i)\}_{i=1}^L$.

4.4 Implementation details

Our diffusion model is a modified SANA DiT [53, 82] adapted for image conditioning, utilizing its pre-trained v1.1 autoencoder. We use UnidepthV2 [54] for metric depth estimation, Mask2Former [13] (trained on Cityscapes [14]) for semantic segmentation, and UniFlow [87] for flow prediction. We trained the model for 2 days on 8 L40 (48GB) GPUs. During training, we drop each condition with a 10% probability. We filter the radar point clouds to a ± 50 m range and use a 512×512 grid. Further implementation details are in the Supplementary Material.

5 Experiments

We start by describing the dataset, metrics, and baseline. In Sec. 5.1, we evaluate *RadarGen* on radar point cloud generation from images. We then show how our proposed method can be used for scene editing in Sec. 5.2. Finally, in Sec. 5.3 we analyze our design choices.

Dataset. We evaluate on the MAN TruckScenes dataset [23] of driving scenes, which provides multi-view camera images and radar point cloud data with higher

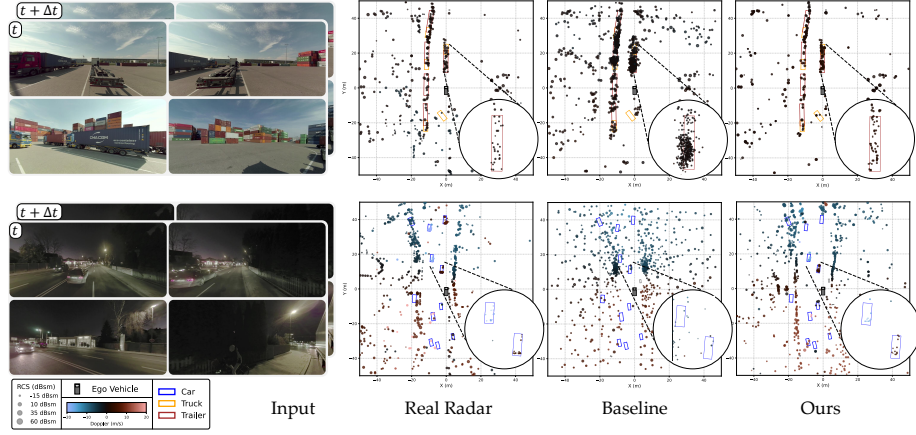


Fig. 5: Qualitative results. Our model generates point clouds with higher geometric and attribute fidelity to the ground truth compared to the baseline. *RadarGen* uses inputs t and $t + \Delta t$, while the baseline uses only t . Ground truth bounding boxes are highlighted in color.

quality than other datasets, and includes rural and urban environments. After excluding dark, uninformative scenes with near-zero visibility, we use 431 clips for training and 49 for evaluation. While clips contain approximately 200 frames each, only 40 include bounding box annotations. We train on all frames in the training set and evaluate only on the annotated frames, which include bounding boxes for visible objects. Evaluation on the nuScenes dataset [8] is available in the Supplementary Material.

Evaluation metrics. As no standard benchmarks exist for our conditional radar generation task, we propose an evaluation framework. Our metrics assess three key aspects: geometric fidelity, radar attribute fidelity, and distribution similarity. Given the importance of foreground points for object detection, we calculate metrics for both the overall scene and points within annotated object bounding boxes. For geometric fidelity, we report *Chamfer Distance* on location (CD Loc.) and the full vector of normalized location, RCS and Doppler (CD Full), Point Cloud *IoU* at 1m (IoU@1m) [63], *Density Similarity*, which measures the distance in number of points in a bounding box compared to the ground truth, and Bounding Box *Hit Rate*. Radar attribute fidelity is assessed using *Distance-Attribute* (DA) Recall, Precision, and F1-score, which assess the proximity of generated points with similar RCS and Doppler values to ground truth points. Finally, we evaluate distribution similarity using *Maximum Mean Discrepancy* (MMD) for both entire point clouds and for points aggregated within object bounding boxes across scenes. Additional details on all metrics are available in the Supplementary Material.

Baseline. We use the feedforward model RGB2Point [39], which is suitable for the task of predicting a point cloud from multi-view images. We extend

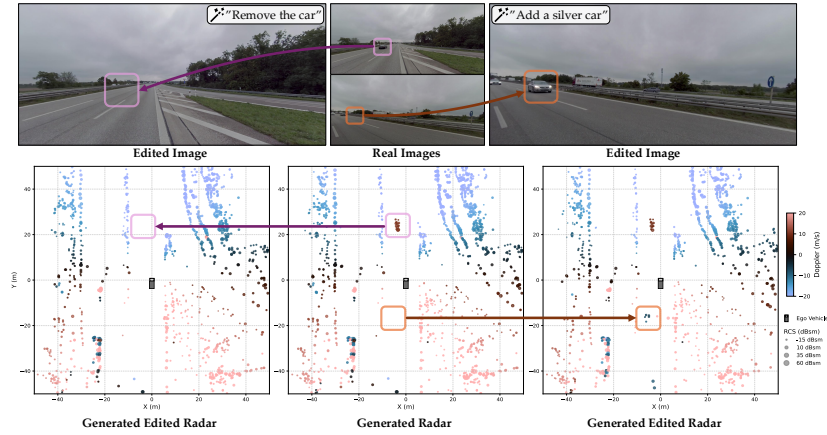


Fig. 6: Scene editing. Modifying the input images using an off-the-shelf image editing tool updates the radar response, demonstrating object removal (left) and insertion (right).

the model output to include RCS and Doppler, and increase the number of parameters to 432M, which is comparable to our model with 592M parameters. See Supplementary Material for more implementation details about this baseline.

5.1 Radar point cloud generation

Quantitative evaluation. In Tab. 1 we compare *RadarGen* to the baseline on the MAN TruckScenes dataset. *RadarGen* broadly outperforms the baseline. A notable exception is CD Full in Entire Area, which is expected as the baseline model was trained using a similar loss objective.

Qualitative evaluation. In Fig. 5, we visually compare our results with the baseline and ground truth on two examples. *RadarGen*’s generated point clouds closely match the ground truth in shape, distribution, and count, demonstrating a significant advantage over the baseline. Further visualizations and scene videos are available in the Supplementary Material.

Compatibility with perception models. We adapt and train the Voxel-NeXt [12] detector on MAN TruckScenes radar data, which yields an NDS [8] of 0.48 on real data (50m range). We then evaluate this detector on generated radar point clouds (PCLs). On PCLs from *RadarGen*, the detector scored an NDS of 0.30. In contrast, the detector struggled to find valid objects in the baseline’s PCLs, resulting in an NDS of nearly zero. Full results are available in the Supplementary Material. Interestingly, while our generated radar achieves a high hit rate (0.66) within true bounding boxes, the overall detection quality still underperforms compared to real data. This could be attributed to the detector’s tailoring to specific, intricate properties of the true data that our model may

Table 2: Ablation Study. We demonstrate the importance of each *RadarGen* condition and compare against a model conditioned directly on multi-view (MV) camera images. Evaluation covers geometric fidelity (CD, IoU, Density Similarity, Hit Rate), radar attribute fidelity (DA Recall, Precision, F1), and distribution similarity (MMD). Shown is the mean $\pm$ s.d., computed over all timestamps for the Entire Area and over all foreground objects for the Foreground.

Method	Entire Area									
	CD Loc. (↓)	CD Full (↓)	IoU @ 1m (↑)	DA Recall (↑)	DA Prec. (↑)	DA F1 (↑)	MMD Loc. (↓)	MMD RCS (↓)	MMD Doppler (↓)	
MV Camera Cond.	1.88 $\pm$ 0.50	0.041 $\pm$ 0.009	0.28 $\pm$ 0.12	0.26 $\pm$ 0.13	0.25 $\pm$ 0.12	0.25 $\pm$ 0.11	0.059 $\pm$ 0.057	0.06 $\pm$ 0.09	0.24 $\pm$ 0.85	
W/o Appearance map	1.71 $\pm$ 0.40	0.040 $\pm$ 0.008	0.31 $\pm$ 0.11	0.23 $\pm$ 0.12	0.25 $\pm$ 0.12	0.23 $\pm$ 0.12	0.059 $\pm$ 0.069	0.09 $\pm$ 0.16	0.35 $\pm$ 0.86	
W/o Semantic map	1.72 $\pm$ 0.40	0.041 $\pm$ 0.010	0.31 $\pm$ 0.11	0.22 $\pm$ 0.12	0.24 $\pm$ 0.12	0.23 $\pm$ 0.12	0.059 $\pm$ 0.061	0.12 $\pm$ 0.26	0.33 $\pm$ 0.72	
W/o Velocity map	1.69 $\pm$ 0.40	0.040 $\pm$ 0.008	0.31 $\pm$ 0.11	0.23 $\pm$ 0.12	0.25 $\pm$ 0.12	0.23 $\pm$ 0.12	0.057 $\pm$ 0.064	0.09 $\pm$ 0.16	0.34 $\pm$ 0.80	
RadarGen	1.68 $\pm$ 0.39	0.040 $\pm$ 0.008	0.31 $\pm$ 0.11	0.23 $\pm$ 0.12	0.26 $\pm$ 0.12	0.24 $\pm$ 0.12	0.056 $\pm$ 0.062	0.09 $\pm$ 0.15	0.31 $\pm$ 0.74	

Method	Foreground									
	CD Loc. (↓)	CD Full (↓)	Density Sim. (↑)	Hit Rate (↑)	MMD Car (↓)		MMD Truck (↓)		MMD Trailer (↓)	
					Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler	Loc. RCS Doppler
MV Camera Cond.	1.0 $\pm$ 0.67	0.069 $\pm$ 0.050	0.47 $\pm$ 0.42	0.56	0.025 0.018	0.024	0.017 0.070	0.073	0.0029 0.037	0.040
W/o Appearance map	0.95 $\pm$ 0.64	0.069 $\pm$ 0.050	0.51 $\pm$ 0.41	0.65	0.034 0.006	0.019	0.019 0.019	0.059	0.0048 0.026	0.041
W/o Semantic map	0.96 $\pm$ 0.66	0.070 $\pm$ 0.050	0.50 $\pm$ 0.41	0.64	0.045 0.010	0.017	0.023 0.039	0.078	0.0072 0.032	0.061
W/o Velocity map	0.95 $\pm$ 0.65	0.069 $\pm$ 0.049	0.51 $\pm$ 0.41	0.66	0.033 0.006	0.019	0.024 0.018	0.070	0.0051 0.024	0.047
RadarGen	0.95 $\pm$ 0.65	0.069 $\pm$ 0.049	0.51 $\pm$ 0.41	0.66	0.037 0.006	0.014	0.024 0.031	0.060	0.0069 0.022	0.046

not fully capture. A detailed analysis of the subtle differences between real and generated radar data is beyond the scope of this paper.

Figure 1 visualizes detections on our generated PCL and demonstrates an augmentation scenario. In this example, a real vehicle is replaced with a generated truck; the detector successfully perceives the newly generated points as a truck.

5.2 Scene editing

Our method supports radar point cloud augmentation by editing the input images using off-the-shelf image editing tools, such as ChronoEdit [78]. Fig. 1 demonstrates object replacement, while Fig. 6 shows examples of object removal and insertion. Notably, in the Fig. 1 replacement example (car to truck), the model correctly removes radar points from the area newly occluded by the truck, demonstrating that it properly handles occlusion changes.

5.3 Ablation study and analysis

BEV conditioning ablations. We assess the contribution of each BEV condition by zeroing it out. Results are shown in Tab. 2. Removing the segmentation map causes the most significant degradation. It worsens geometric fidelity and significantly increases the RCS MMD for both the Entire Area and per-object class in the Foreground. Ablating either the velocity map or the appearance map primarily degrades the Doppler MMD. Notably, the degradation from removing the appearance map, even with the segmentation map present, suggests the model leverages finer-grained appearance details to refine its understanding of object class and thus generate a more realistic motion profile. This confirms that all input channels are meaningfully fused to produce radar point clouds.

Comparison to multi-view conditioning. We also compare against a model conditioned directly on multi-view (MV) camera images, which omits the BEV

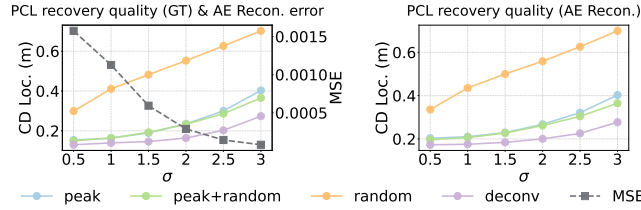


Fig. 7: Impact of σ and recovery method on PCL reconstruction. **peak** selects local maxima; **random** samples pixels proportional to the probability map (without replacement); **peak+random** takes peaks first, then fills remaining points by probability sampling excluding the peaks. MSE is calculated between the input and output of the AE, whose range is $[-1,1]$.

input. Instead, it uses images from t and $t + \Delta t$ concatenated as input tokens with Plücker and modality embeddings. While this MV model achieves an improved MMD for radar attributes on the Entire Area, our BEV conditioned model yields better overall geometric fidelity, as detailed in Tab. 2. Furthermore, the MV approach is computationally prohibitive, requiring over $3\times$ runtime; this model was trained for 9 days, compared to only 2 days for our BEV conditioned model. While our BEV representation provides a more efficient and geometrically accurate solution, the results from the MV model highlight it as a valuable direction for future research.

Radar PCL reconstruction. We ablate two factors: (a) the 2D Gaussian kernel bandwidth σ for the Point Density Map, and (b) the PCL recovery method. We test $\sigma \in \{0.5, 1, 1.5, 2, 2.5, 3\}$ and four recovery methods (**random**, **peak**, **peak+random**, **deconv**) on the MAN TruckScenes *mini-train* set, evaluating AE reconstruction error and the geometric fidelity of the recovered PCL. As shown in Fig. 7 (left), larger σ values reduce AE reconstruction error, as smoother maps are easier to reconstruct. However, overly large σ over-smooths the map structure, degrading downstream PCL recovery. Balancing this trade-off, we set $\sigma = 2$. For the recovery method, **deconv** consistently yields the best PCL quality across all σ values and for both ground-truth and AE-reconstructed maps, validating its ability to preserve the spatial distribution.

6 Conclusion and Future Work

In this paper, we introduced *RadarGen*, a probabilistic diffusion framework for generating realistic automotive radar point clouds from multi-view camera inputs. Our method leverages foundation models to create a unified BEV representation for conditioning, generates dense BEV radar maps, and uses a deconvolution solver to recover the final sparse point cloud. Experimentally, *RadarGen* outperforms our proposed baseline on a comprehensive suite of geometric, attribute, and distribution metrics. We demonstrate its application in data augmentation via simple image editing and show promising results for downstream

detectors. Future work will focus on extending our framework for video input, exploring text-based conditioning, and training on multiple datasets and radar configurations.

Limitations. Our method’s performance is inherently tied to the capabilities of the upstream foundation models. It is thus limited in challenging scenarios where these models underperform, such as in low-light night scenes, with strong reflections, or during camera occlusion. Additionally, our model can generate points in areas not directly visible to the cameras. This behavior is a double-edged sword: while it is desirable for "filling in" occluded objects, it can also lead to uncontrolled hallucinations in these unseen regions. Finally, in this work we did not model explicit radar mechanisms, as doing so requires large-scale raw data and detailed radar specifications and internal mechanisms, which are often undisclosed. If large-scale raw data becomes available, modeling explicit radar properties would be a fruitful next step.

Acknowledgments

Or Litany acknowledges support from the Israel Science Foundation (grant 624/25) and the Azrieli Foundation Early Career Faculty Fellowship. The authors gratefully acknowledge this support. This research was supported by the Council for Higher Education in Israel under the Moonshot Project.

References

1. Achlioptas, P., Diamanti, O., Mitliagkas, I., Guibas, L.: Learning representations and generative models for 3d point clouds. In: ICML (2018)
2. Alkanat, T., Pandharipande, A.: Automotive radar point cloud parametric density estimation using camera images. In: ICASSP (2024)
3. Arnold, M., Bauhofer, M., Mandelli, S., Henninger, M., Schaich, F., Wild, T., ten Brink, S.: Maxray: A raytracing-based integrated sensing and communication framework. In: IEEE JC&S (2022)
4. Beck, A., Teboulle, M.: A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM journal on imaging sciences* (2009)
5. Bialer, O., Haitman, Y.: Radsimreal: Bridging the gap between synthetic and real data in radar object detection with simulation. In: CVPR (2024)
6. Borts, D., Liang, E., Broedermann, T., Ramazzina, A., Walz, S., Palladin, E., Sun, J., Brueggemann, D., Sakaridis, C., Van Gool, L., et al.: Radar fields: Frequency-space neural scene representations for fmcw radar. In: ACM SIGGRAPH (2024)
7. Caccia, L., Van Hoof, H., Courville, A., Pineau, J.: Deep generative modeling of lidar data. In: IROS (2019)
8. Caesar, H., Bankiti, V., Lang, A.H., Vora, S., Liong, V.E., Xu, Q., Krishnan, A., Pan, Y., Baldan, G., Beijbom, O.: nuscenes: A multimodal dataset for autonomous driving. In: CVPR (2020)
9. Cai, R., Yang, G., Averbuch-Elor, H., Hao, Z., Belongie, S., Snavely, N., Hariharan, B.: Learning gradient fields for shape generation. In: ECCV (2020)
10. Capsoni, C., D’Amico, M.: A physically based radar simulator. *Journal of Atmospheric and Oceanic Technology* (1998)

11. Chen, X., Zhang, X.: Rf genesis: Zero-shot generalization of mmwave sensing through simulation-based data synthesis and generative diffusion models. In: Proceedings of the 21st ACM Conference on Embedded Networked Sensor Systems (2023)
12. Chen, Y., Liu, J., Zhang, X., Qi, X., Jia, J.: Voxelnex: Fully sparse voxelnet for 3d object detection and tracking. In: CVPR (2023)
13. Cheng, B., Misra, I., Schwing, A.G., Kirillov, A., Girdhar, R.: Masked-attention mask transformer for universal image segmentation. In: CVPR (2022)
14. Cordts, M., Omran, M., Ramos, S., Rehfeld, T., Enzweiler, M., Benenson, R., Franke, U., Roth, S., Schiele, B.: The cityscapes dataset for semantic urban scene understanding. In: CVPR (2016)
15. Daubechies, I., DeVore, R., Fornasier, M., Güntürk, C.S.: Iteratively reweighted least squares minimization for sparse recovery. *Communications on Pure and Applied Mathematics: A Journal Issued by the Courant Institute of Mathematical Sciences* (2010)
16. De Oliveira, M.L.L., Bekooij, M.J.: Generating synthetic short-range fmcw range-doppler maps using generative adversarial networks and deep convolutional autoencoders. In: IEEE Radar Conference (2020)
17. Deng, K., Zhao, D., Han, Q., Zhang, Z., Wang, S., Zhou, A., Ma, H.: Midas: Generating mmwave radar data from videos for training pervasive and privacy-preserving human sensing tasks. *Proceedings of the ACM on IMWUT* (2023)
18. Ding, F., Palffy, A., Gavrilu, D.M., Lu, C.X.: Hidden gems: 4d radar scene flow learning using cross-modal supervision. In: CVPR (2023)
19. Ding, F., Pan, Z., Deng, Y., Deng, J., Lu, C.X.: Self-supervised scene flow estimation with 4-d automotive radar. *IEEE RA-L* (2022)
20. Ding, F., Wen, X., Zhu, Y., Li, Y., Lu, C.X.: Radarocc: Robust 3d occupancy prediction with 4d imaging radar. *NeurIPS* (2024)
21. Dosovitskiy, A., Ros, G., Codevilla, F., Lopez, A., Koltun, V.: Carla: An open urban driving simulator. In: Conference on robot learning (2017)
22. Dudek, M., Kissinger, D., Weigel, R., Fischer, G.: A millimeter-wave fmcw radar system simulator for automotive applications including nonlinear component models. In: EuRAD (2011)
23. Fent, F., Kutteneich, F., Ruch, F., Rizwin, F., Juergens, S., Lechermann, L., Nissler, C., Perl, A., Voll, U., Yan, M., et al.: Man truckscenes: A multimodal dataset for autonomous trucking in diverse conditions. *NeurIPS* (2024)
24. Fidelis, E.C., Reway, F., Ribeiro, H., Campos, P.L., Huber, W., Icking, C., Faria, L.A., Schön, T.: Generation of realistic synthetic raw radar data for automated driving applications using generative adversarial networks. *arXiv preprint arXiv:2308.02632* (2023)
25. Giannakis, I., Giannopoulos, A., Warren, C.: A realistic fdtd numerical modeling framework of ground penetrating radar for landmine detection. *IEEE journal of selected topics in applied earth observations and remote sensing* (2015)
26. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., Bengio, Y.: Generative adversarial networks. *Communications of the ACM* (2020)
27. Gubelli, D., Krasnov, O.A., Yaroyvi, O.: Ray-tracing simulator for radar signals propagation in radar networks. In: EuRAD (2013)
28. He, D., Guan, K., Ai, B., Zhong, Z., Kim, J., Chung, H., Hrovat, A.: Channel measurement and ray-tracing simulation for 77 ghz automotive radar. *IEEE TITS* (2022)

29. He, K., Liang, R., Munkberg, J., Hasselgren, J., Vijaykumar, N., Keller, A., Fidler, S., Gilitschenski, I., Gojcic, Z., Wang, Z.: Unirelight: Learning joint decomposition and synthesis for video relighting. arXiv preprint arXiv:2506.15673 (2025)
30. Hirsenkorn, N., Subkowski, P., Hanke, T., Schaermann, A., Rauch, A., Rasshofer, R., Biebl, E.: A ray launching approach for modeling an fmcw radar system. In: IRS (2017)
31. Ho, J., Jain, A., Abbeel, P.: Denoising diffusion probabilistic models. arXiv preprint arXiv:2006.11239 (2020)
32. Holder, M., Linnhoff, C., Rosenberger, P., Winner, H.: The fourier tracing approach for modeling automotive radar sensors. In: IRS (2019)
33. Hu, Q., Zhang, Z., Hu, W.: Rangeldm: Fast realistic lidar point cloud generation. In: ECCV (2024)
34. Huang, T., Miller, J., Prabhakara, A., Jin, T., Laroia, T., Kolter, Z., Rowe, A.: Dart: Implicit doppler tomography for radar novel view synthesis. In: CVPR (2024)
35. Ke, B., Obukhov, A., Huang, S., Metzger, N., Daudt, R.C., Schindler, K.: Repurposing diffusion-based image generators for monocular depth estimation. In: CVPR (2024)
36. Kerbl, B., Kopanas, G., Leimkühler, T., Drettakis, G.: 3d gaussian splatting for real-time radiance field rendering. ACM Trans. Graph. (2023)
37. Kingma, D.P., Welling, M., et al.: An introduction to variational autoencoders. Foundations and Trends® in Machine Learning (2019)
38. Kung, P.C., Harisha, S., Vasudevan, R., Eid, A., Skinner, K.A.: Radarsplat: Radar gaussian splatting for high-fidelity data synthesis and 3d reconstruction of autonomous driving scenes. arXiv preprint arXiv:2506.01379 (2025)
39. Lee, J.J., Benes, B.: Rgb2point: 3d point cloud generation from single rgb images. In: WACV (2025)
40. Lei, Z., Xu, F., Wei, J., Cai, F., Wang, F., Jin, Y.Q.: Sar-nerf: Neural radiance fields for synthetic aperture radar multi-view representation. IEEE TGRS (2024)
41. Li, A., Lei, Z., Wei, J., Xu, F.: Sar-gs: 3d gaussian splatting for synthetic aperture radar target reconstruction. arXiv preprint arXiv:2506.21633 (2025)
42. Lin, C.H., Wang, Z., Liang, R., Zhang, Y., Fidler, S., Wang, S., Gojcic, Z.: Controllable weather synthesis and removal with video diffusion models. IEEE/CVF International Conference on Computer Vision (ICCV) (2025)
43. Liu, H., Xing, B., Wang, H., Cui, J., Spencer, B.F.: Simulation of ground penetrating radar on dispersive media by a finite element time domain algorithm. Journal of applied geophysics (2019)
44. Luo, S., Hu, W.: Diffusion probabilistic models for 3d point cloud generation. In: CVPR (2021)
45. Melas-Kyriazi, L., Rupprecht, C., Vedaldi, A.: Pc2: Projection-conditioned point cloud diffusion for single-image 3d reconstruction. In: CVPR (2023)
46. Mildenhall, B., Srinivasan, P.P., Tancik, M., Barron, J.T., Ramamoorthi, R., Ng, R.: Nerf: Representing scenes as neural radiance fields for view synthesis. Communications of the ACM (2021)
47. Mo, K., Guerrero, P., Yi, L., Su, H., Wonka, P., Mitra, N.J., Guibas, L.J.: StructureNet: hierarchical graph networks for 3d shape generation. ACM TOG (2019)
48. Nakashima, K., Kurazume, R.: Lidar data synthesis with denoising diffusion probabilistic models. In: ICRA (2024)
49. Nichol, A., Jun, H., Dhariwal, P., Mishkin, P., Chen, M.: Point-e: A system for generating 3d point clouds from complex prompts. arXiv preprint arXiv:2212.08751 (2022)

50. NVIDIA, Agarwal, N., Ali, A., Bala, M., Balaji, Y., Barker, E., Cai, T., Chattopadhyay, P., Chen, Y., Cui, Y., Ding, Y., Dworakowski, D., Fan, J., Fenzi, M., Ferroni, F., Fidler, S., Fox, D., Ge, S., Ge, Y., Gu, J., Gururani, S., He, E., Huang, J., Huffman, J., Jannaty, P., Jin, J., Kim, S.W., Klár, G., Lam, G., Lan, S., Leal-Taixe, L., Li, A., Li, Z., Lin, C.H., Lin, T.Y., Ling, H., Liu, M.Y., Liu, X., Luo, A., Ma, Q., Mao, H., Mo, K., Mousavian, A., Nah, S., Niverty, S., Page, D., Paschalidou, D., Patel, Z., Pavao, L., Ramezanali, M., Reda, F., Ren, X., Sabavat, V.R.N., Schmerling, E., Shi, S., Stefaniak, B., Tang, S., Tchapmi, L., Tredak, P., Tseng, W.C., Varghese, J., Wang, H., Wang, H., Wang, H., Wang, T.C., Wei, F., Wei, X., Wu, J.Z., Xu, J., Yang, W., Yen-Chen, L., Zeng, X., Zeng, Y., Zhang, J., Zhang, Q., Zhang, Y., Zhao, Q., Zolkowski, A.: Cosmos world foundation model platform for physical ai (2025), <https://arxiv.org/abs/2501.03575>
51. Ouza, M., Ulrich, M., Yang, B.: A simple radar simulation tool for 3d objects based on blender. In: IRS (2017)
52. Pan, Z., Ding, F., Zhong, H., Lu, C.X.: Ratrack: moving object detection and tracking with 4d radar point cloud. In: ICRA (2024)
53. Peebles, W., Xie, S.: Scalable diffusion models with transformers. In: ICCV (2023)
54. Piccinelli, L., Sakaridis, C., Yang, Y.H., Segu, M., Li, S., Abbeloos, W., Gool, L.V.: UniDepthV2: Universal monocular metric depth estimation made simpler (2025), <https://arxiv.org/abs/2502.20110>
55. Rafidashti, M., Lan, J., Fatemi, M., Fu, J., Hammarstrand, L., Svensson, L.: Neuradar: Neural radiance fields for automotive radar point clouds. In: CVPR (2025)
56. Ran, H., Guizilini, V., Wang, Y.: Towards realistic scene generation with lidar diffusion models. In: CVPR (2024)
57. Rangaraj, P.A., Alkanat, T., Pandharipande, A.: Raids: Radar range-azimuth map estimation from image, depth, and semantic descriptions. IEEE Sensors Journal (2025)
58. Remcom Inc.: Wavefarer® automotive radar software (2024)
59. Ren, X., Lu, Y., Cao, T., Gao, R., Huang, S., Sabour, A., Shen, T., Pfaff, T., Wu, J.Z., Chen, R., et al.: Cosmos-drive-dreams: Scalable synthetic driving data generation with world foundation models. arXiv preprint arXiv:2506.09042 (2025)
60. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models (2021)
61. Russell, L., Hu, A., Bertoni, L., Fedoseev, G., Shotton, J., Arani, E., Corrado, G.: Gaia-2: A controllable multi-view generative world model for autonomous driving. arXiv preprint arXiv:2503.20523 (2025)
62. Sallab, A.E., Sobh, I., Zahran, M., Essam, N.: Lidar sensor modeling and data augmentation with gans for autonomous driving. arXiv preprint arXiv:1905.07290 (2019)
63. Saunders, K., Manso, L.J., Vogiatzis, G.: Baseboostdepth: Exploiting larger baselines for self-supervised monocular depth estimation. arXiv preprint arXiv:2407.20437 (2024)
64. Scharf, L.L., Demeure, C.: Statistical signal processing: detection, estimation, and time series analysis. Prentice Hall (1991)
65. Schöffmann, C., Ubezio, B., Böhm, C., Mühlbacher-Karrer, S., Zangl, H.: Virtual radar: Real-time millimeter-wave radar sensor simulation for perception-driven robotics. IEEE RA-L (2021)
66. Schüßler, C., Hoffmann, M., Bräunig, J., Ullmann, I., Ebel, R., Vossiek, M.: A realistic radar ray tracing simulator for large mimo-arrays in automotive environments. IEEE Journal of Microwaves (2021)

67. Sligar, A.P.: Machine learning-based radar perception for autonomous vehicles using full physics simulation. *IEEE Access* (2020)
68. Song, P., Song, D., Yang, Y., Lan, E., Liu, J.: Simulating automotive radar with lidar and camera inputs. *arXiv preprint arXiv:2503.08068* (2025)
69. Stetco, C., Ubezio, B., Mühlbacher-Karrer, S., Zangl, H.: Radar sensors in collaborative robotics: Fast simulation and experimental validation. In: *IEEE ICRA* (2020)
70. Stowell, M.L., Fasnfest, B.J., White, D.A.: Investigation of radar propagation in buildings: A 10-billion element cartesian-mesh ftd simulation. *IEEE transactions on antennas and propagation* (2008)
71. Thieling, J., Frese, S., Roßmann, J.: Scalable and physical radar sensor simulation for interacting digital twins. *IEEE Sensors Journal* (2020)
72. Tibshirani, R.: Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society Series B: Statistical Methodology* (1996)
73. Topak, A.E., Hasch, J., Zwick, T.: A system simulation of a 77 ghz phased array radar sensor. In: *IRS* (2011)
74. Tyszkiewicz, M.J., Fua, P., Trulls, E.: Gecco: Geometrically-conditioned point diffusion models. In: *ICCV* (2023)
75. Vahdat, A., Williams, F., Gojcic, Z., Litany, O., Fidler, S., Kreis, K., et al.: Lion: Latent point diffusion models for 3d shape generation. *NeurIPS* (2022)
76. Weston, R., Jones, O.P., Posner, I.: There and back again: Learning to simulate radar data for real-world applications. In: *IEEE ICRA* (2021)
77. Wheeler, T.A., Holder, M., Winner, H., Kochenderfer, M.J.: Deep stochastic radar models. In: *IEEE IV* (2017)
78. Wu, J.Z., Ren, X., Shen, T., Cao, T., He, K., Lu, Y., Gao, R., Xie, E., Lan, S., Alvarez, J.M., Gao, J., Fidler, S., Wang, Z., Ling, H.: Chronoedit: Towards temporal reasoning for image editing and world simulation. *arXiv preprint arXiv:2510.04290* (2025)
79. Wu, Y., Zhang, K., Qian, J., Xie, J., Yang, J.: Text2lidar: Text-guided lidar point cloud generation via equirectangular transformer. In: *ECCV* (2024)
80. Wu, Z., Wang, Y., Feng, M., Xie, H., Mian, A.: Sketch and text guided diffusion model for colored point cloud generation. In: *ICCV* (2023)
81. Xiao, W., Huang, H., Zhong, C., Lin, Y., Wang, N., Chen, X., Chen, Z., Zhang, S., Yang, S., Meriaux, P., et al.: Simulate any radar: Attribute-controllable radar simulation via waveform parameter embedding. *arXiv preprint arXiv:2506.03134* (2025)
82. Xie, E., Chen, J., Chen, J., Cai, H., Tang, H., Lin, Y., Zhang, Z., Li, M., Zhu, L., Lu, Y., et al.: Sana: Efficient high-resolution image synthesis with linear diffusion transformers. *arXiv preprint arXiv:2410.10629* (2024)
83. Xiong, Y., Ma, W.C., Wang, J., Urtasun, R.: Learning compact representations for lidar completion and generation. In: *CVPR* (2023)
84. Yang, G., Huang, X., Hao, Z., Liu, M.Y., Belongie, S., Hariharan, B.: Pointflow: 3d point cloud generation with continuous normalizing flows. In: *ICCV* (2019)
85. Yun, Z., Iskander, M.F.: Ray tracing for radio propagation modeling: Principles and applications. *IEEE access* (2015)
86. Zamorski, M., Zięba, M., Klukowski, P., Nowak, R., Kurach, K., Stokowiec, W., Trzciniński, T.: Adversarial autoencoders for compact representations of 3d point clouds. *Computer Vision and Image Understanding* (2020)
87. Zhang, Y., Keetha, N., Lyu, C., Jhamb, B., Chen, Y., Qiu, Y., Karhade, J., Jha, S., Hu, Y., Ramanan, D., Scherer, S., Wang, W.: Ufm: A simple path towards unified dense correspondence with flow. In: *arXiv* (2025)

- 88. Zhou, L., Du, Y., Wu, J.: 3d shape generation and completion through point-voxel diffusion. In: ICCV (2021)
- 89. Zyrianov, V., Zhu, X., Wang, S.: Learning to generate realistic lidar point clouds. In: ECCV (2022)