

Robust ISAC Object Tracking via Cross-Modal Supervision and Spatio-Temporal Skip-Transformer

Won Jae Ryu , *Member, IEEE*, Sanjay Bhardwaj , Jae-Min Lee , *Member, IEEE*,
and Dong-Seong Kim , *Senior Member, IEEE*

Abstract—Robust object tracking in ISAC systems is challenging due to the inherent sparsity and severe impulsive noise of millimeter-wave (mmWave) radar signals. To address this, we propose a novel tracking framework that combines cross-modal supervision with a Spatio-Temporal Skip-Transformer (ST-Skipformer). An offline LiDAR–camera label generator produces high-fidelity ground-truth supervision, enabling the ST-Skipformer to learn precise spatial features from noisy radar inputs. Furthermore, the ST-Skipformer incorporates a Temporal Skip (TS) Block to effectively filter background clutter while recovering high-frequency spatial details via skip connections. Experimental results on the real-world DeepSense 6 G dataset demonstrate that the proposed method achieves an Average Displacement Error (ADE) of 0.0174 m, reducing localization error by approximately 98% compared to baselines, while attaining a near-perfect F1-score of 0.9967.

Index Terms—Integrated sensing and communication (ISAC), mmWave object tracking, cross-modal learning, spatio-temporal transformer, radar signal processing.

I. INTRODUCTION

INTEGRATED Sensing and Communication (ISAC) is a cornerstone 6 G technology, with recent advancements expanding into diverse domains, including Reconfigurable Intelligent Surface (RIS)-empowered satellite Internet of Things (IoT) for bridging coverage-efficiency gaps in sensing and access [2]. However, a significant gap remains between theoretical advancements and practical terrestrial deployment due to the heavy reliance on synthetic datasets [1], [3], [4]. These datasets often fail to capture real-world radio channel complexities like multi-path fading and impulsive noise. While recent real-world datasets have emerged, they present distinct limitations: automotive datasets like nuScenes [5] provide only processed radar point clouds lacking raw signal features, whereas existing studies

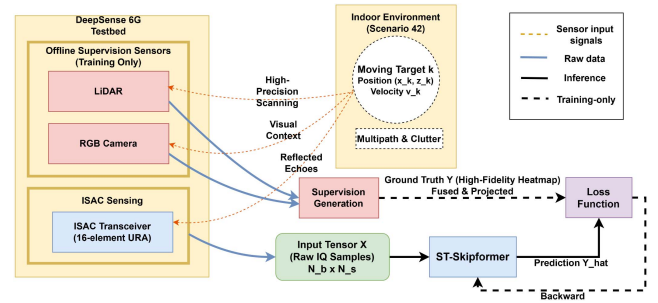


Fig. 1. Overview of the proposed framework.

utilizing raw datasets like DeepSense 6G [6] have predominantly focused on communication tasks (e.g., beam prediction) [7], [8], leaving precise object localization under-explored. A critical bottleneck is the difficulty of obtaining high-fidelity ground truth for raw radio-frequency (RF) signals, which are unintuitive for manual annotation [9]. To bridge this gap, we propose a novel framework leveraging cross-modal supervision for robust object tracking. Unlike previous works relying on noisy camera-radar fusion [9], we utilize a co-located high-precision Light Detection and Ranging (LiDAR) sensor to generate high-fidelity ground truth heatmaps. While optical and LiDAR sensors may degrade in extreme weather, our cross-modal label generator is explicitly designed to operate entirely offline in controlled indoor or clear-weather environments, transferring high-fidelity spatial precision to the radar-based ST-Skipformer. Our contributions are: 1) the Spatio-Temporal Skip-Transformer (ST-Skipformer), a dense spatio-temporal architecture with Temporal Skip (TS) Blocks for robust localization on raw ISAC signals; 2) a cross-modal supervision pipeline fusing LiDAR and camera to generate high-fidelity ground-truth heatmaps for radar-based learning; and 3) validation on real-world datasets demonstrating a 98% error reduction compared to four representative baselines spanning dense convolutional (U-Net), recurrent (CRNN), attention-based (CNN-ViT), and graph-based (RadarGNN) paradigms. The remainder of this paper is organized as follows. Section II describes the system model, Section III details the ST-Skipformer, and Section IV presents the results.

II. SYSTEM MODEL

We consider a downlink ISAC scenario in an indoor environment, as illustrated in Fig. 1. The experimental setup is based

Received 6 May 2026; accepted 3 June 2026. Date of publication 8 June 2026; date of current version 30 June 2026. This work was supported in part by the Institute for Information and Communications Technology Planning and Evaluation (IITP) under Grant IITP-2026-RS-2020-II201612 (25%) and Grant IITP-2026-RS-2024-00438430 (25%), and in part by the National Research Foundation of Korea (NRF), through the Ministry of Education under Grant 2018R1A6A1A03024003 (25%) and Grant RS-2025-25431637 (25%). The associate editor coordinating the review of this article and approving it for publication was Prof. Li You. (Corresponding author: Dong-Seong Kim.)

Won Jae Ryu and Sanjay Bhardwaj are with the ICT Convergence Research Center, Kumoh National Institute of Technology, Gumi 39177, Republic of Korea (e-mail: wj0828@kumoh.ac.kr; sanjay.b1976@gmail.com).

Jae-Min Lee and Dong-Seong Kim are with the Department of IT Convergence Engineering, Kumoh National Institute of Technology, Gumi 39177, Republic of Korea (e-mail: ljmpaul@kumoh.ac.kr; dskim@kumoh.ac.kr).

Digital Object Identifier 10.1109/LSP.2026.3701218

on the DeepSense 6 G dataset (Scenario 42) [6], [10]. While the original testbed provides various sensing modalities, we specifically utilize the synchronized data from the 60 GHz mmWave ISAC transceiver, the mechanical LiDAR, and the RGB camera. As depicted in Fig. 1, the ISAC transceiver captures raw RF signals reflected from moving targets and is the sole sensor used at inference time. The co-located high-precision LiDAR and RGB camera are used only during offline training: by fusing the semantic information from the camera with the precise depth measurements from the LiDAR, high-fidelity ground-truth heatmaps $\mathbf{Y}$ are produced as cross-modal supervision labels for training the ISAC network.

A. ISAC Signal Model

The ISAC system operates at $f_c = 60$ GHz with $B = 1$ GHz bandwidth [10]. The transmitter employs a 16-element (2×8) Uniform Rectangular Array (URA) and performs sequential beam sweeping using a codebook of $N_b = 21$ beams. Unlike simultaneous multi-beam systems, the transceiver illuminates each direction $n \in \{1, \dots, N_b\}$ sequentially within a single sweep cycle. The received signal for the n -th beam is modeled as:

$$y_n(t) = \sum_{k=1}^K \alpha_{n,k} e^{j2\pi f_{D,k}t} x_n(t - \tau_k) + w(t), \quad (1)$$

where n , $\alpha_{n,k}$, $f_{D,k}$, and τ_k represent the temporal order of beam activation, the channel gain, Doppler shift, and delay, respectively, K denotes the number of targets in the scene, and $w(t)$ represents additive white Gaussian noise (AWGN).

B. Problem Formulation

The objective is to estimate the 2D spatial distribution of targets using only the ISAC signal $\mathbf{X}$ during inference. Let $Y \in [0, 1]^{H \times W}$ denote the ground truth occupancy heatmap. To ensure both semantic correctness and spatial precision, Y is derived by fusing the object detection results from the RGB camera with the high-precision depth measurements from the LiDAR sensor. The collected raw signals are preprocessed into a multi-channel spectrogram tensor $\mathbf{X} \in \mathbb{R}^{T \times C \times H \times W}$, where T is the number of temporal frames, $C = N_b$ is the number of beam-direction channels, and $H \times W$ is the spatial resolution of each per-beam spectrogram. Formally, we aim to learn a non-linear mapping function $\mathcal{F}_\theta : \mathbb{R}^{T \times C \times H \times W} \rightarrow [0, 1]^{H \times W}$, parameterized by θ , to estimate the spatial distribution of targets. The training objective is to find the optimal parameters θ^* that minimize the discrepancy between the prediction $\hat{\mathbf{Y}} = \mathcal{F}_\theta(\mathbf{X})$ and the ground truth $\mathbf{Y}$. This can be formulated as the following optimization problem:

$$\theta^* = \arg \min_{\theta} \mathbb{E}_{(\mathbf{X}, \mathbf{Y}) \sim \mathcal{D}} [\mathcal{L}(\mathcal{F}_\theta(\mathbf{X}), \mathbf{Y})], \quad (2)$$

where $\mathcal{D}$ represents the training dataset derived from the DeepSense 6G scenarios, and $\mathcal{L}(\cdot, \cdot)$ denotes the loss function designed to handle the sparsity and imbalance of the heatmap labels.

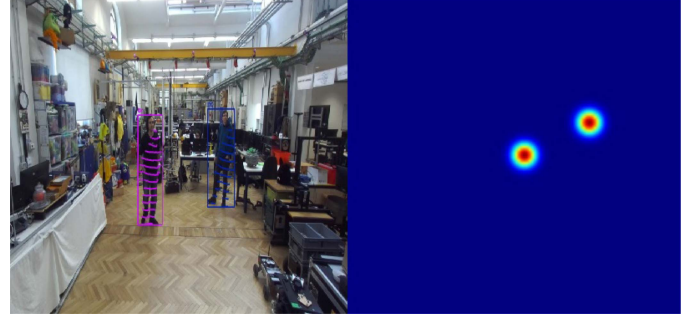


Fig. 2. The cross-modal supervision generation.

III. PROPOSED METHOD

To achieve robust object tracking using raw ISAC signals, we propose a cross-modal supervision framework as illustrated in Fig. 1. During offline training, a LiDAR–camera pipeline generates high-fidelity ground-truth heatmaps, which supervise the ST-Skipformer—the radar-only network that performs inference at deployment time.

A. Cross-Modal Supervision Generation

We exploit a co-located 32-channel LiDAR and RGB camera to generate high-fidelity labels, leveraging the LiDAR's 2 cm range resolution [10]. After static background removal (distance threshold 0.1 m), 3D LiDAR points are projected onto the camera image plane via the extrinsic $[R|t]$ and intrinsic K matrices, yielding a valid pixel-aligned point set $\mathcal{P}_{\text{valid}}$. For each YOLOv8-detected human bounding box B_i , we extract the sub-cloud $\mathcal{P}_i = \{p \in \mathcal{P}_{\text{valid}} \mid \text{proj}(p) \in B_i\}$ and obtain a robust 3D centroid via depth median. YOLOv8 was selected for its strong person-detection accuracy, mature pretrained release, and real-time throughput; the detector is modular and any comparable 2D detector can be substituted without affecting downstream training. The centroids are tracked on a Bird's-Eye View (BEV) plane using a linear Kalman Filter with state $[x, z, v_x, v_z]^T$ and Hungarian data association, then mapped onto a 128×128 grid. The ground-truth heatmap $\mathbf{Y}$ is generated by a $\sigma=3.0$ 2D Gaussian:

$$\mathbf{Y}(i, j) = \max_{k \in \mathcal{O}} \exp \left(-\frac{(i - x_k)^2 + (j - z_k)^2}{2\sigma^2} \right). \quad (3)$$

The qualitative result is shown in Fig. 2.

B. ISAC Signal Preprocessing

To extract spatial features from the time-sequential beam data, we transform each beam's in-phase and quadrature (IQ) samples into a 2D Range-Time Map (RTM) via the Short-Time Fourier Transform (STFT). A single ISAC frame at time t is then constructed by RTMs from a complete sweep cycle of N_b beams. The resulting multi-channel spectrogram at time t is resized and normalized to form a spatial feature map $x_t \in \mathbb{R}^{C \times 128 \times 128}$, where the channel dimension $C = N_b$ corresponds to the individual beam directions.

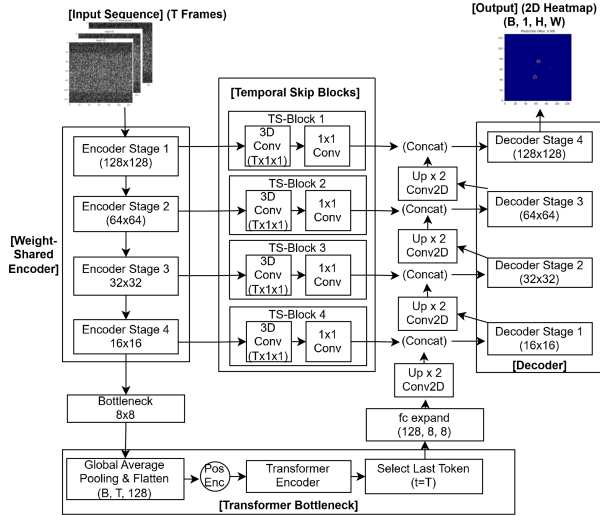


Fig. 3. ST-Skipformer architecture.

C. ST-Skipformer Architecture

To effectively handle the inherent noise in radar ISAC data while maintaining temporal consistency, we propose the ST-Skipformer, whose overall architecture is illustrated in Fig. 3. Unlike the standard U-Net [14], which directly transfers encoder features to the decoder, our architecture introduces a novel TS-Block that filters impulsive noise using 3D convolutions.

1) *Weight-Shared Encoder*: The network takes a sequence of T preprocessed ISAC frames as input, denoted as $\mathbf{X} \in \mathbb{R}^{B \times T \times C \times H \times W}$. Here, B represents the batch size, T is the number of temporal frames, C denotes the channel dimension representing the spatial beams ($C = N_b$), and $H \times W$ refers to the spatial resolution of the spectrogram (e.g., 128×128). To extract spatial features from each frame independently, we employ a weight-shared CNN encoder, as shown in the left part of Fig. 3. We reshape the input tensor to $(B \cdot T, C, H, W)$ and pass it through four convolutional blocks. Let $f_\theta(\cdot)$ denote the encoder function; the feature map extraction for the t -th frame at the l -th layer is given by:

$$\mathbf{F}_t^{(l)} = f_\theta^{(l)}(\mathbf{x}_t), \quad t = 1, \dots, T \quad (4)$$

where $\mathbf{F}_t^{(l)}$ represents the spatial feature map of the t -th frame.

2) *Temporal Skip Connection*: Standard skip connections directly transfer encoder features to the decoder, which often propagates high-frequency artifacts and impulsive noise in ISAC signals, leading to overfitting. Particularly in sequential beamforming systems, moving targets can introduce phase inconsistencies across different beam directions within a single sweep. To mitigate this, we design the TS-Block as shown in Fig. 3. The TS-Block aggregates the feature sequence $\mathcal{F}^{(l)} = \{\mathbf{F}_1^{(l)}, \dots, \mathbf{F}_T^{(l)}\}$ into a single, noise-suppressed feature map $\mathbf{S}^{(l)}$. We stack the sequence along the temporal dimension to form a 5D tensor $\mathbf{V}^{(l)} \in \mathbb{R}^{B \times C \times T \times H_l \times W_l}$. A 3D Convolution with a kernel size of $(T, 1, 1)$ is then applied to compress the entire temporal axis:

$$\mathbf{Z}^{(l)} = \sigma \left(\text{GroupNorm} \left(\text{Conv3d}_{T \times 1 \times 1}(\mathbf{V}^{(l)}) \right) \right), \quad (5)$$

where σ denotes the ReLU activation. This operation effectively learns to amplify temporally consistent signals (targets) while suppressing transient noise. Finally, a 1×1 2D convolution is applied to reduce channel dimensions before concatenation with the decoder features. Based on our empirical analysis of the trained 3D convolutional kernels, the network learns nearly uniform temporal weights. This indicates that the TS-Block effectively functions as a Learned Non-Coherent Integrator. By performing temporal averaging, it effectively suppresses uncorrelated transient noise and multi-path clutter while maximizing the Signal-to-Noise Ratio (SNR) of consistent moving targets.

3) *Transformer Bottleneck*: To capture long-range temporal dependencies, we employ a Transformer at the bottleneck, depicted at the bottom of Fig. 3. The spatial features of the deepest layer $\mathbf{F}_t^{(4)}$ are compressed into a latent vector via Global Average Pooling (GAP):

$$\mathbf{e}_t = \text{GAP}(\mathbf{F}_t^{(4)}) \in \mathbb{R}^{d_{\text{model}}}. \quad (6)$$

The sequence of embeddings $\mathbf{E} = [\mathbf{e}_1, \dots, \mathbf{e}_T]$ is augmented with learnable positional encodings and processed by a Transformer encoder consisting of $L = 2$ stacked layers, each comprising 4-head self-attention and a feed-forward network with hidden dimension $2 \cdot d_{\text{model}}$, where $d_{\text{model}} = 128$. This multi-layer self-attention captures long-range temporal dependencies that simple pooling cannot model. The final context vector $\mathbf{c}_{\text{final}}$, representing the state at the last timestamp, is expanded by a fully-connected layer and reshaped into a $128 \times 8 \times 8$ tensor, which initiates the decoder upsampling cascade.

4) *Decoder and Hybrid Loss Function*: The decoder reconstructs the spatial resolution through four upsampling stages. At each stage, a $2 \times$ transposed convolution doubles the spatial resolution of the decoder feature, which is then concatenated with the matching-resolution TS-Block output $\mathbf{S}^{(l)}$ before a 3×3 convolution refines the merged tensor. To train the model, we utilize a hybrid loss function combining Gaussian Focal Loss (L_{focal}) and Soft Dice Loss (L_{dice}) to address the class imbalance between the sparse targets and the background. The Gaussian Focal Loss [11] focuses on the peaks of the heatmap:

$$L_{\text{focal}} = \frac{-1}{N} \sum_{xy} \begin{cases} (1 - \hat{Y}_{xy})^\alpha \log(\hat{Y}_{xy}), & \text{if } Y_{xy} = 1, \\ (1 - Y_{xy})^\beta (\hat{Y}_{xy})^\alpha \log(1 - \hat{Y}_{xy}), & \text{otherwise,} \end{cases} \quad (7)$$

where $\hat{Y}$ is the predicted heatmap, Y is the ground truth, and N is the number of objects.

The Soft Dice Loss [12] maximizes the overlap between the prediction and the target distribution:

$$L_{\text{dice}} = 1 - \frac{2 \sum \hat{Y} \cdot Y + \epsilon}{\sum \hat{Y} + \sum Y + \epsilon}. \quad (8)$$

The total objective function is defined as $L_{\text{total}} = L_{\text{focal}} + L_{\text{dice}}$.

IV. EXPERIMENTAL RESULTS

We evaluate the ST-Skipformer using Scenario 42 of the DeepSense 6 G Dataset [6], [10]. This scenario involves human targets with low Radar Cross Section (RCS), causing severe

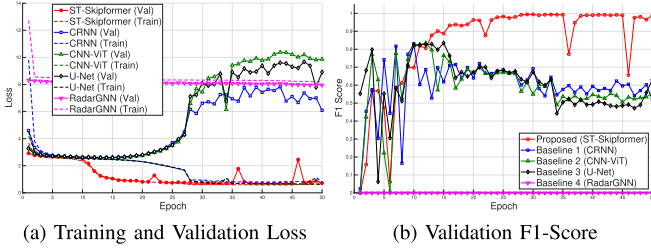


Fig. 4. Training progress of ST-Skipformer and baselines.

 TABLE I
PERFORMANCE COMPARISON ON THE TEST SET

Model	ADE ↓	F1-Score ↑	Precision ↑	Recall ↑
U-Net	0.8486	0.7766	0.8067	0.7488
CRNN	0.8887	0.8146	0.8010	0.8287
CNN-ViT	0.9967	0.8017	0.8188	0.7852
RadarGNN	12.0	0.0	0.0	0.0
ST-Skipformer	0.0174	0.9967	0.9983	0.9951

signal fluctuations. The dataset (2,100 frames) was partitioned into training, validation, and testing sets (70:15:15). Training was performed for 50 epochs using AdamW [13] ($lr = 10^{-3}$) on an NVIDIA RTX 4060 GPU. We compare our model against Frame-Stacked U-Net [14], CRNN [15], CNN-ViT [16], and RadarGNN [17], where RadarGNN represents a peak-based graph neural network paradigm originally designed for pre-processed automotive radar point clouds. Fig. 4 illustrates the training dynamics. As shown in Fig. 4(a), ST-Skipformer converges rapidly within 15 epochs and remains stable. In contrast, baseline models exhibit significant generalization gaps, overfitting to impulsive noise and background clutter. Fig. 4(b) further proves the efficiency of the proposed cross-modal strategy as our F1-score quickly reaches saturation. The mild oscillations of the validation loss in the later epochs reflect statistical variance over the limited validation pool (~ 315 frames) and are not accompanied by any degradation of the validation F1-score, which saturates near 0.998. This indicates that the fluctuations stem from sample-level loss variance rather than overfitting. Table I summarizes the test results. ST-Skipformer achieves an ADE of 0.0174 m, a 98% reduction compared to the best baseline (CRNN, 0.8887 m). This validates that the TS-Block effectively suppresses impulsive noise, recovering spatial details lost in standard architectures. With a near-perfect precision of 0.9983 (only 1 false positive among 605 predictions on the test set), our model effectively eliminates ghost targets caused by multipath reflections, whereas baselines show significant performance limitations.

Fig. 5 reveals a striking pattern: U-Net, CRNN, and CNN-ViT all produce diffuse responses concentrated near the spatial centroid of the training-set target distribution rather than at the actual targets, indicating that they have memorized the marginal position prior rather than learning the input-to-location mapping. RadarGNN collapses entirely. In contrast, ST-Skipformer produces sharp responses co-located with the ground-truth targets, confirming that it has learned a true functional mapping from

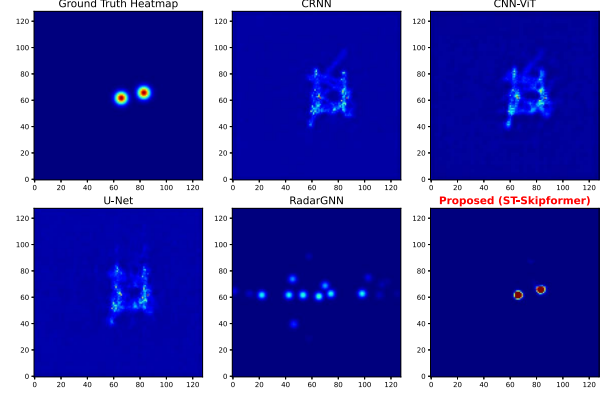


Fig. 5. Qualitative comparison of object tracking heatmaps between the proposed ST-Skipformer and baseline models.

the noisy RTM tensor to spatial positions and that the TS-Blocks effectively suppress impulsive noise.

Among the baselines, RadarGNN exhibits complete failure ($F1 = 0$, $ADE = 12$ m). This stems from a paradigm mismatch: RadarGNN's peak-based graph construction requires stable scatterer peaks, but under low-RCS human returns in the ISAC RTM, target reflections are buried beneath speckle noise, so the extracted point cloud is effectively random. This empirically reinforces the necessity of dense-tensor architectures over sparsifying paradigms for raw ISAC perception. The full ST-Skipformer contains 1.95 M parameters and 1.97 G FLOPs per inference, of which the TS-Blocks account for only 240 K parameters (12.3% of total). This is enabled by the (T, 1, 1) temporal kernel design, which avoids the cost of a generic 3D convolution.

V. CONCLUSION

This letter presented a robust ISAC object tracking framework designed to overcome the limitations of sparse and noisy mmWave signals. By establishing a cross-modal supervision pipeline that leverages fused LiDAR-camera labels, we successfully bridged the modality gap and provided precise offline supervision for radar-based learning. Crucially, the proposed ST-Skipformer with TS-Blocks effectively suppressed impulsive noise and preserved local spatial information often lost in conventional Transformers. Extensive evaluations on real-world data confirmed the superiority of our approach, achieving an ADE of 0.0174 m and an F1-score of 0.9967, significantly outperforming U-Net, CRNN, CNN-ViT, and RadarGNN baselines, including those representing graph-based radar perception paradigms. Although our evaluation is confined to Scenario 42, the fact that competing architectures collapse to mean-response predictions even within this scenario shows that learning a true RTM-to-position mapping is itself a non-trivial prerequisite for multi-scenario validation. Future work will explore multi-modal robust fusion techniques for the offline label generation pipeline to maintain high-quality supervision even under severe weather conditions, and extend the evaluation across additional DeepSense 6 G scenarios and outdoor benchmarks to further validate its generalizability.

REFERENCES

- [1] J. Wang, N. Varshney, C. Gentile, S. Blandino, J. Chuang, and N. Golmie, "Integrated sensing and communication: Enabling techniques, applications, tools and data sets, standardization, and future directions," *IEEE Internet Things J.*, vol. 9, no. 23, pp. 23416–23440, Dec. 2022.
- [2] Y. Xie, Z. Lin, R. Ma, K. An, X. Zhong, and Y. He, "RIS-Empowered satellite IoT: Bridging the coverage-efficiency gap of last-mile access and sensing," *IEEE Internet Things Mag.*, early access, Feb. 9, 2026, doi: [10.1109/MIOT.2026.3658612](https://doi.org/10.1109/MIOT.2026.3658612).
- [3] S. Lu et al., "Integrated sensing and communications: Recent advances and ten open challenges," *IEEE Internet Things J.*, vol. 11, no. 11, pp. 19094–19120, Jun. 2024.
- [4] J. A. Zhang et al., "Enabling joint communication and radar sensing in mobile networks—A survey," *IEEE Commun. Surveys Tuts.*, vol. 24, no. 1, pp. 306–345, 1st Quart. 2022.
- [5] H. Caesar et al., "nuScenes: A multimodal dataset for autonomous driving," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.*, 2020, pp. 11621–11631.
- [6] A. Alkhateeb et al., "DeepSense 6G: A large-scale real-world multi-modal sensing and communication dataset," *IEEE Commun. Mag.*, vol. 61, no. 9, pp. 122–128, Sep. 2023.
- [7] U. Demirhan and A. Alkhateeb, "Radar aided 6G beam prediction: Deep learning algorithms and real-world demonstration," in *Proc. IEEE Wireless Commun. Netw. Conf.*, Austin, TX, USA, 2022, pp. 2655–2660.
- [8] G. Charan, T. Osman, A. Hredzak, N. Thawdar, and A. Alkhateeb, "Vision-position multi-modal beam prediction using real millimeter wave datasets," in *Proc. IEEE Wireless Commun. Netw. Conf.*, Austin, TX, USA, 2022, pp. 2727–2731.
- [9] Y. Wang, Z. Jiang, X. Gao, J.-N. Hwang, G. Xing, and H. Liu, "RODNet: Radar object detection using cross-modal supervision," in *Proc. IEEE Winter Conf. Appl. Comput. Vis.*, Jan. 2021, pp. 504–513.
- [10] Wireless Intelligence Lab, "DeepSense 6G dataset: Scenarios 42-44," 2024. [Online]. Available: <https://deepsense6g.net/scenarios/Scenarios%2040-49/scenarios-42-44>
- [11] H. Law and J. Deng, "CornerNet: Detecting objects as paired keypoints," in *Proc. Eur. Conf. Comput. Vis.*, 2018, pp. 734–750.
- [12] F. Milletari, N. Navab, and S.-A. Ahmadi, "V-Net: Fully convolutional neural networks for volumetric medical image segmentation," in *Proc. 4th Int. Conf. 3D Vis.*, Stanford, CA, USA, 2016, pp. 565–571.
- [13] I. Loshchilov and F. Hutter, "Decoupled weight decay regularization," in *Proc. Int. Conf. Learn. Representations*, 2017.
- [14] O. Ronneberger, P. Fischer, and T. Brox, "U-Net: Convolutional networks for biomedical image segmentation," in *Proc. 18th Int. Conf. Med. Image Comput. Comput.-Assist. Intervent.*, 2015, pp. 234–241.
- [15] B. Shi, X. Bai, and C. Yao, "An end-to-end trainable neural network for image-based sequence recognition and its application to scene text recognition," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 39, no. 11, pp. 2298–2304, Nov. 2017.
- [16] H. Wu et al., "CvT: Introducing convolutions to vision transformers," in *Proc. IEEE Int. Conf. Comput. Vis.*, 2021, pp. 22–31.
- [17] F. Fent, P. Bauerschmidt, and M. Lienkamp, "RadarGNN: Transformation invariant graph neural network for radar-based perception," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. Workshops*, 2023, pp. 182–191.