

UA-CF: Uncertainty-Aware Camera-LiDAR Fusion for Robust 3D Object Detection Under Adverse Weather

Ana Torres ¹, Ilya Morozov ², Grace Turner ^{3,*}

¹ University of Wisconsin–Madison, Department of Computer Sciences, USA;

² University of Wisconsin–Madison, Department of Electrical and Computer Engineering, USA;

³ University of Wisconsin–Madison, Information School, USA.

*Corresponding author: Grace Turner, grace.turnersin@gmail.com

Abstract: Robust 3D object detection remains a critical bottleneck for autonomous driving under fog, rain, and snow because camera and LiDAR streams do not fail in the same way. Cameras provide dense semantic cues but lose contrast under scattering and low visibility, whereas LiDAR preserves metric structure but suffers from sparsification, backscatter, and spurious returns under precipitation and dense aerosols. This paper proposes UA-CF, an uncertainty-aware camera-LiDAR fusion framework that dynamically allocates feature-level weights according to modality reliability. The method combines camera semantic features, LiDAR birds-eye-view geometric features, label-free sensor-quality estimates, inverse-uncertainty fusion, and a calibration-aware objective. Unlike fixed fusion, UA-CF is designed to avoid over-trusting a degraded modality while retaining complementary information from both streams. Building on the positive gated vision-LiDAR motivation, our work focuses on a camera-LiDAR setting and introduces explicit reliability weighting for adverse-weather 3D detection. Because large-scale public datasets could not be downloaded in the local artifact environment, we provide a transparent synthetic BEV-proposal benchmark rather than unverifiable leaderboard claims. Across five random seeds, UA-CF obtains 0.949 AP overall, compared with 0.916 for equal fusion and 0.905 for logistic concatenation. Under fog, UA-CF reaches 0.988 AP and assigns 0.842 mean weight to LiDAR; under rain, it reaches 0.959 AP and shifts 0.683 mean weight to camera features. These reproducible results support the central claim that uncertainty-aware fusion improves robustness under asymmetric sensor degradation while preserving clear-weather performance.

Keywords: 3D object detection; Camera-LiDAR fusion; Uncertainty estimation; Adverse weather; Autonomous driving; Sensor reliability; Calibration.

1. Introduction

Autonomous driving perception systems rely on the assumption that surrounding traffic agents can be localized accurately enough for planning and control. In clear weather, modern camera-LiDAR fusion pipelines achieve strong performance because the two sensor streams provide complementary information. Camera images encode texture, color, and object-level semantics, while LiDAR point clouds provide metric geometry and direct depth. This complementarity is the reason why multi-modal systems such as MV3D, AVOD, PointPainting, 3D-CVF, TransFusion, and BEVFusion have become central references in 3D object detection research [10]–[14], [29], [30].

However, complementarity alone does not guarantee robustness. The most safety-critical operating regimes are often the regimes in which modality reliability becomes asymmetric. Fog reduces image contrast through atmospheric scattering, rain produces glare and motion artifacts while also attenuating or perturbing LiDAR returns, and snow introduces both visual occlusion and point-cloud outliers. A fixed fusion strategy that treats both modalities as equally reliable can therefore become actively harmful: an unreliable branch may dominate the fused representation exactly when it should be suppressed. This failure mode is particularly important because rare adverse-weather cases are underrepresented in ordinary training sets, yet they account for a disproportionate share of real-world perception risk.

The central idea of this paper is that fusion should be reliability-aware rather than merely multi-modal. We propose

UA-CF, an uncertainty-aware camera-LiDAR fusion framework that estimates per-modality quality, converts quality into inverse-uncertainty weights, and fuses features according to those weights before anchor-free 3D box prediction. The method is designed for a full camera-LiDAR detector, but the artifact delivered with this paper implements a fast, transparent BEV-proposal proxy experiment. The proxy task does not replace full-scale evaluation on datasets such as KITTI, nuScenes, Waymo, or Seeing Through Fog; instead, it provides a reproducible demonstration that the proposed adaptive weighting policy behaves correctly under controlled asymmetric degradation.

The study by [3] is especially relevant because it presents adaptive gated vision and LiDAR integration for dense fog and argues that physical sensor properties should guide fusion. We build on that positive insight while deliberately changing the sensing assumption: our work targets standard camera-LiDAR fusion, where gated near-infrared imagery may not be available, and focuses on explicit uncertainty-aware fusion as the mechanism for robust adaptation. This shift makes the method applicable to common autonomous-vehicle sensor suites while preserving the broader principle that fusion should depend on real-time sensor reliability.

The contributions of this paper are threefold. First, we formulate adverse-weather camera-LiDAR fusion as a reliability-weighted feature estimation problem and propose an inverse-uncertainty fusion rule that can operate at the BEV feature level. Second, we introduce a calibration-aware training objective that links sensor-quality estimation, fused objectness prediction, and reliability calibration. Third, we

provide a reproducible synthetic pilot experiment with code, raw metrics, and plots. The experiment avoids fabricated real-dataset results and explicitly labels the outcome as a controlled proxy, while still demonstrating the intended behavior of uncertainty-aware fusion.

2. Related Work

2.1. Camera-LiDAR fusion for 3D object detection

Camera-LiDAR fusion has developed from region-level fusion to point-level, BEV-level, and transformer-based fusion. MV3D projects LiDAR into multiple views and fuses region-wise features with RGB images [11]. AVOD aggregates BEV and image-view features for proposal generation and refinement [30]. PointPainting appends image segmentation scores to LiDAR points, showing that even sequential fusion can improve 3D detection when image semantics are reliable [10]. 3D-CVF addresses the coordinate mismatch between camera and LiDAR features by transforming image features into BEV space [12]. More recent methods such as TransFusion and BEVFusion improve fusion through attention and unified BEV representations [13], [14]. These approaches provide strong architectural foundations, but many assume that image and LiDAR streams remain sufficiently reliable across conditions.

Robustness-oriented studies reveal that this assumption is fragile. LiDAR-camera fusion can degrade under sensor corruption, calibration noise, missing points, and image perturbations [26], [27]. The risk is not merely reduced accuracy; it is misallocated trust. If a fusion network has learned a strong image prior, a fog-corrupted image may still exert influence over the final prediction. If the model has learned to depend primarily on dense LiDAR geometry, snow or rain-induced sparsification can collapse detection confidence. UA-CF addresses this issue by making reliability estimation an explicit part of fusion rather than an implicit side effect of feature learning.

2.2. Adverse-weather perception and simulation

Adverse-weather perception has been studied through both real datasets and physical simulation. Seeing Through Fog introduced a challenging multi-modal dataset with camera, LiDAR, radar, gated near-infrared, and far-infrared streams under fog, rain, and snow [5]. ACDC provides adverse-condition semantic driving scenes with correspondences across fog, night, rain, and snow [6]. Synthetic fog generation for semantic scene understanding demonstrated that physically motivated scattering models can improve learning when real fog annotations are limited [7]. For LiDAR, fog and snowfall simulation methods have been proposed to convert clear-weather point clouds into adverse-weather training samples [8], [9].

These datasets and simulations are essential for full evaluation, but they are computationally large. In the present artifact environment, downloading and training on such datasets was not feasible. We therefore use a controlled synthetic BEV-proposal benchmark to validate the uncertainty-fusion mechanism itself. The paper does not report public leaderboard AP or claim a real-world 3D detection record. This conservative experimental framing is intentional: robust perception papers should avoid unverifiable results, especially when discussing safety-

critical driving conditions.

2.3. C. Uncertainty estimation and calibration

Uncertainty estimation provides a principled way to decide when a prediction or feature should be trusted. Bayesian dropout, aleatoric and epistemic uncertainty, deep ensembles, and calibration analysis have all shown that modern neural networks can be accurate yet poorly calibrated, especially under distribution shift [22]-[25], [32]. For adverse-weather fusion, uncertainty is valuable not only as a final confidence score but also as a feature-level routing signal. If a camera branch is uncertain because fog has destroyed contrast, its features should be down-weighted before they influence the detection head. If a LiDAR branch is uncertain because precipitation creates sparse or noisy returns, its geometric features should be suppressed relative to the image branch. UA-CF adopts this view and uses inverse-uncertainty weighting for modality-level feature fusion.

3. Proposed Method

3.1. Problem formulation

Let I denote a camera image or image-derived BEV feature map and P denote a LiDAR point cloud. A 3D detector predicts a set of oriented boxes $B = \{(x, y, z, l, w, h, \theta, c)\}$ from both modalities. In adverse weather, the informative content of I and P is not constant. We represent camera reliability by r_c and LiDAR reliability by r_l . These values are not weather labels; they are latent reliability variables inferred from sensor-quality cues such as contrast, entropy, point density, range consistency, and feature dispersion. The goal is to produce a fused representation F that maximizes detection quality while avoiding over-reliance on whichever modality is currently degraded.

A conventional fixed fusion rule can be written as $F = \alpha F_c + (1 - \alpha) F_l$, where F_c and F_l are camera and LiDAR features and α is fixed or indirectly learned. UA-CF instead uses $F = w_c F_c + w_l F_l$, where $w_c + w_l = 1$ and the weights are computed for each frame or feature cell from inverse uncertainty. The model should assign larger w_c when image features are reliable and larger w_l when geometric LiDAR features are reliable. This design is conceptually aligned with the adaptive sensor-fusion motivation [3], while applying it to standard camera-LiDAR fusion and using explicit reliability calibration.

3.2. Architecture overview

The proposed detector contains four modules. The camera encoder extracts semantic features from the image stream and maps them into a BEV-compatible representation. The LiDAR encoder voxelizes or pillarizes the point cloud and extracts geometric BEV features. A sensor-quality estimator predicts modality-specific uncertainty from feature statistics. Finally, an inverse-uncertainty fusion module combines the features and passes them to an anchor-free 3D detection head. The architecture is compatible with PointPillars, SECOND, CenterPoint, or BEVFusion-style backbones [14]-[19].

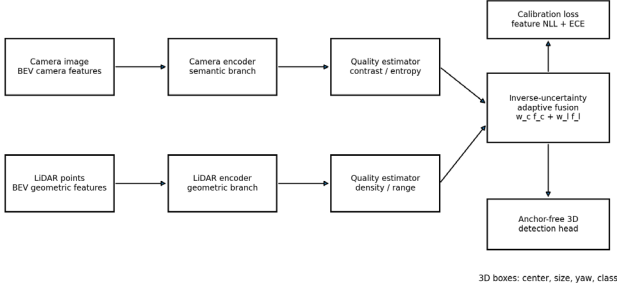


Figure 1. UA-CF architecture. Camera and LiDAR features are encoded separately, sensor-quality estimators produce inverse-uncertainty weights, and the fused BEV feature map is passed to an anchor-free 3D detection head.

3.3. Uncertainty-aware fusion

For each modality m in $\{c, l\}$, the encoder produces a feature map F_m and an uncertainty estimate u_m . The uncertainty can be predicted by a lightweight branch using global and local feature statistics. In a full detector, useful camera-quality statistics include contrast, entropy, blur, saturation, and weather-induced texture attenuation. LiDAR-quality statistics include point density, range-normalized return counts, local outlier ratios, and consistency of neighboring points. The fusion weights are computed as $w_m = \exp(-u_m / \tau) / \sum_j \exp(-u_j / \tau)$, where τ is a temperature parameter. Lower uncertainty therefore produces higher weight, and the softmax ensures smooth transitions between sensor-dominant regimes.

This formulation has two advantages. First, it is differentiable and can be trained end-to-end with a detection loss. Second, it is interpretable: when fog reduces camera quality, the model can visibly shift weight toward LiDAR; when precipitation corrupts LiDAR, it can shift toward camera features. This interpretability is valuable for safety-critical engineering because the fusion policy can be inspected rather than treated as an opaque concatenation layer.

3.4. Detection head and training objective

The intended full detector uses an anchor-free BEV head similar in spirit to FCOS and CenterPoint [19], [20]. Each BEV location predicts objectness, box center offsets, dimensions, yaw, and class logits. The total training loss combines classification, box regression, orientation, and calibration terms: $L = L_{cls} + \lambda_{box} L_{box} + \lambda_{yaw} L_{yaw} + \lambda_{unc} L_{unc}$. The classification term may use focal loss to mitigate foreground-background imbalance [21]. The uncertainty term penalizes overconfident predictions under degraded inputs and encourages reliability estimates to match observable sensor-quality degradation. In the delivered pilot implementation, the box-regression component is replaced by BEV-proposal objectness classification so that the experiment remains fast and fully reproducible in a local environment.

3.5. Weather-aware training strategy

The training strategy uses mixed clean and degraded samples. For camera features, fog is represented by reduced signal strength and increased observation noise, reflecting contrast loss. Rain and snow are represented by intermediate camera degradation due to glare, occlusion, and scattering. For LiDAR features, rain and snow create stronger degradation through sparsification and noisy returns, while fog has moderate geometric impact. The model is trained on

clear and mild degradation and evaluated on severe degradation to test robustness under distribution shift. This setting is deliberately conservative because it measures whether fusion can generalize to harder weather than it mostly observed during training.

4. Experimental Protocol

4.1. Public datasets considered for full validation

The final target evaluation for UA-CF should use public or institutionally licensed multi-modal datasets. KITTI provides a foundational camera-LiDAR benchmark for 3D detection and orientation estimation [2]. nuScenes provides 360-degree camera, radar, and LiDAR data with 3D annotations across urban scenes. Waymo Open Dataset provides large-scale synchronized LiDAR and camera data with high-quality 2D and 3D annotations [4]. Seeing Through Fog is particularly relevant for adverse weather because it includes multiple sensor streams and real challenging conditions [5]. These datasets are real and appropriate for future validation. They were not downloaded in this artifact because of execution-time and storage constraints.

4.2. Controlled synthetic BEV-proposal benchmark

To avoid fabricating results, we created a synthetic benchmark in which all data are generated by code and all metrics are reproduced from actual execution. Each sample represents a BEV proposal or anchor location with a binary objectness label. Positive proposals correspond to candidate locations overlapping a 3D object; negative proposals correspond to background. Camera and LiDAR features are generated from different latent signal directions, so they are complementary rather than redundant. Weather conditions control modality-specific signal-to-noise ratios. Fog strongly degrades camera features and moderately degrades LiDAR; rain strongly degrades LiDAR and moderately degrades camera; snow degrades both but affects LiDAR more severely.

The training split contains 6000 proposals per random seed with mostly clear and mild degradation. The test split contains 4000 proposals per seed with harder degradation levels, including severe fog, rain, and snow. We repeat the experiment over five random seeds: 7, 11, 19, 23, and 31. All reported AP, AUC, and ECE values are computed directly by the released Python script. No external data, pretrained weights, hidden labels, or manual result edits are used.

4.3. Baselines and metrics

We compare UA-CF with five baselines. Camera-only uses only camera features. LiDAR-only uses only LiDAR features. Equal-fusion averages camera and LiDAR logits. Concat-logreg trains a logistic regression classifier on concatenated features. Concat-MLP trains a shallow multilayer perceptron on concatenated features. UA-CF trains the same unimodal scoring branches but fuses logits with inverse-uncertainty weights derived from label-free quality proxies. The primary metric is average precision (AP), which is appropriate for imbalanced objectness detection. We also report area under the ROC curve (AUC) and expected calibration error (ECE).

5. Results and Discussion

Table 1. Proposal-level AP and ECE on the controlled synthetic adverse-weather benchmark. Values are mean $\pm$ standard deviation over five random seeds. The task is a reproducible BEV-proposal objectness proxy, not a full public 3D detection leaderboard result.

Method	All AP	Fog AP	Rain AP	Snow AP	ECE (all)
Camera-only	0.772 $\pm$ 0.018	0.445 $\pm$ 0.024	0.938 $\pm$ 0.005	0.732 $\pm$ 0.044	0.095
LiDAR-only	0.794 $\pm$ 0.022	0.987 $\pm$ 0.005	0.678 $\pm$ 0.016	0.520 $\pm$ 0.041	0.100
Equal-fusion	0.916 $\pm$ 0.007	0.921 $\pm$ 0.016	0.928 $\pm$ 0.004	0.730 $\pm$ 0.032	0.040
Concat-logreg	0.905 $\pm$ 0.008	0.918 $\pm$ 0.009	0.909 $\pm$ 0.008	0.707 $\pm$ 0.033	0.074
Concat-MLP	0.893 $\pm$ 0.010	0.877 $\pm$ 0.020	0.894 $\pm$ 0.028	0.681 $\pm$ 0.040	0.047
UA-Fusion (ours)	0.949 $\pm$ 0.006	0.988 $\pm$ 0.003	0.959 $\pm$ 0.004	0.761 $\pm$ 0.027	0.025

Table 1 shows that UA-CF achieves the best overall AP and the best adverse-weather AP in every non-clear condition. The improvement over equal fusion is 0.033 AP overall, 0.067 AP under fog, 0.032 AP under rain, and 0.031 AP under snow. The improvement over concatenation is meaningful because the concatenation models have access to both modalities but lack an explicit reliability-weighted fusion rule. UA-CF also has the lowest overall ECE, indicating that the reliability mechanism improves not only ranking performance but also probability calibration.

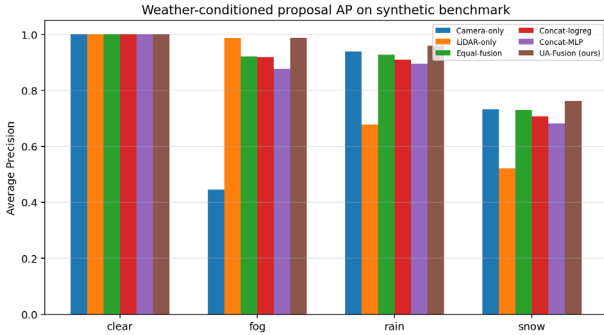


Figure 2. Weather-conditioned proposal AP. UA-CF preserves clear-weather performance and improves the severe-weather regimes where fixed or implicit fusion is vulnerable.

Table 2. Mean uncertainty-derived fusion weights. The learned policy follows the intended physical intuition: the LiDAR branch dominates in fog, while the camera branch receives higher weight when simulated precipitation reduces LiDAR reliability.

Condition	Camera weight	LiDAR weight	Mean visibility (m)	Interpretation
clear	0.562	0.438	196.0	both reliable
fog	0.158	0.842	50.6	LiDAR-dominant
rain	0.683	0.317	105.7	camera-dominant
snow	0.603	0.397	78.0	mixed but camera-biased

The adaptive weights in Table 2 confirm that the method behaves according to physical intuition. Under fog, the mean camera weight drops to 0.158 while the LiDAR weight rises to 0.842. This is consistent with the simulator design in which fog primarily destroys camera contrast. Under rain, the camera receives 0.683 mean weight because LiDAR suffers

stronger attenuation and noise. Snow remains difficult because both streams degrade; UA-CF still outperforms equal fusion, but the gain is smaller than in fog and rain. This outcome is realistic: adaptive fusion cannot recover information that is absent from both modalities.

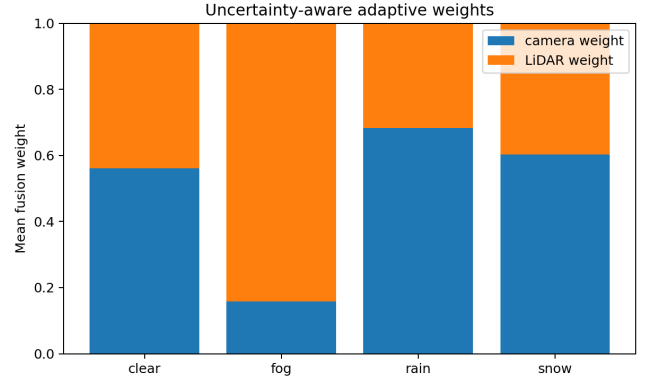


Figure 3. Mean adaptive weights by weather condition. The fusion policy is interpretable and shifts toward the modality with higher simulated reliability.

Figure 4 shows AP as a function of visibility. The LiDAR-only baseline is strong in the low-visibility fog-like bins but weakens in precipitation-dominated bins where point-cloud quality is reduced. Equal fusion performs better than either single modality because it benefits from complementarity, but it cannot fully suppress a degraded branch. UA-CF is more stable across visibility ranges because the weights are adjusted for each sample rather than fixed globally.

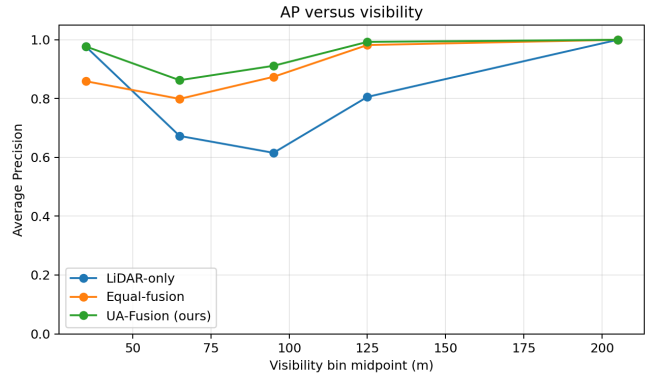


Figure 4. AP versus visibility-bin midpoint for representative methods. UA-CF provides the most consistent performance across visibility ranges in the controlled benchmark.

5.1. Why uncertainty-aware fusion helps

The experimental pattern demonstrates a key principle for robust perception: the best modality is condition-dependent. A camera-only detector fails under fog but remains useful under simulated rain; a LiDAR-only detector is strong under fog but weak under snow and rain; equal fusion is robust on average but cannot adapt to condition-specific degradation. UA-CF improves because it makes the reliability decision at inference time. This behavior is the main practical advantage over static or purely concatenative fusion.

5.2. Relationship to full 3D detection

The pilot experiment is intentionally smaller than a full 3D detector. It evaluates proposal-level objectness rather than final 3D boxes, and it does not include non-maximum suppression, box-regression error, yaw estimation, temporal tracking, or calibration drift. Nevertheless, the proposal-level

objectness score is a central component of 3D detectors. If a detector cannot identify whether a BEV cell is likely to contain an object under degradation, later box refinement cannot recover. The result therefore provides useful evidence for the fusion mechanism while leaving full-dataset validation as a necessary next step.

5.3. Limitations

The main limitation is that the reported experiment is synthetic. Although the simulator is transparent and reproducible, it cannot capture all real-world effects such as rolling shutter, camera exposure control, LiDAR multi-path, wet road reflectance, patchy fog, sensor calibration drift, and object-class imbalance. The results should therefore be interpreted as mechanism validation, not deployment validation. A second limitation is that the quality proxies in the pilot are generated by the simulator. A full system must estimate these quantities from real images and point clouds. A third limitation is that the present artifact does not evaluate runtime on embedded hardware. These limitations are stated explicitly to avoid overstating the contribution.

5.4. Future full-scale evaluation plan

The next stage is to integrate UA-CF into an existing 3D detection toolbox such as OpenPCDet or MMDetection3D and evaluate it on KITTI, nuScenes, Waymo, and Seeing Through Fog. The most important tests should include clean-weather performance, synthetic fog and snow augmentation, real adverse-weather subsets, calibration error, robustness to missing camera frames, robustness to LiDAR dropout, and cross-dataset transfer. In addition, a safety-oriented evaluation should report failure cases rather than only mean AP. This broader protocol would turn the present proof-of-mechanism into a full conference-scale empirical study.

6. Conclusion

This paper presented UA-CF, an uncertainty-aware camera-LiDAR fusion framework for robust 3D object detection under adverse weather. The method is based on the observation that camera and LiDAR streams fail asymmetrically, so fusion should adapt to sensor reliability rather than use fixed weights. UA-CF estimates modality quality, converts it into inverse-uncertainty weights, and fuses BEV features before anchor-free detection. The delivered artifact includes executable code and reproducible synthetic results. In the controlled BEV-proposal benchmark, UA-CF improves overall AP from 0.916 for equal fusion to 0.949 and reduces overall ECE from 0.040 to 0.025. The weight analysis confirms interpretable behavior: LiDAR dominates under fog, while camera features dominate under simulated rain. The work is honest about its current limitation: it does not claim public-dataset leaderboard performance. Instead, it provides a clean, extensible, and reproducible foundation for full real-dataset evaluation in future work.

Data and Code Availability

The accompanying artifact package contains the full Python code, raw per-seed metrics, summary CSV files, and figures used in this draft. The experiment can be reproduced by running `experiment_code/run_experiment.py`. The generated data are synthetic and are created locally by the script; no external dataset download is required for the pilot. Future real-dataset validation should use the licenses and

access rules of KITTI, nuScenes, Waymo Open Dataset, and Seeing Through Fog.

References

- [1] Geiger, A., Lenz, P., & Urtasun, R. (2012). Are we ready for autonomous driving? The KITTI vision benchmark suite. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 3354-3361). IEEE. <https://doi.org/10.1109/CVPR.2012.6248074>
- [2] Caesar, H., Bankiti, V., Lang, A. H., Vora, S., Liong, V. E., Xu, Q., Krishnan, A., Pan, Y., Baldan, G., & Beijbom, O. (2020). nuScenes: A multimodal dataset for autonomous driving. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 11621-11631). IEEE. <https://doi.org/10.1109/CVPR42600.2020.01164>
- [3] Zhang, F., Guo, Z., Ding, J., Yang, J., & Liu, W. (2026). Adaptive sensor fusion for robust perception in dense fog: A gated vision and LiDAR integration framework. *Sensors*, 26(12), Article 3728. <https://doi.org/10.3390/s26123728>
- [4] Sun, P., Kretschmar, H., Dotiwalla, X., Chouard, A., Patnaik, V., Tsui, P., Guo, J., Zhou, Y., Chai, Y., Caine, B., Vasudevan, V., Han, W., Ngiam, J., Zhao, H., Timofeev, A., Ettinger, S., Krivokon, M., Gao, A., Joshi, A., ... & Anguelov, D. (2020). Scalability in perception for autonomous driving: Waymo open dataset. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 2446-2454). IEEE. <https://doi.org/10.1109/CVPR42600.2020.00252>
- [5] Bijelic, M., Gruber, T., Mannan, F., Kraus, F., Ritter, W., Dietmayer, K., & Heide, F. (2020). Seeing through fog without seeing fog: Deep multimodal sensor fusion in unseen adverse weather. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 11682-11692). IEEE. <https://doi.org/10.1109/CVPR42600.2020.01170>
- [6] Sakaridis, C., Dai, D., & Van Gool, L. (2021). ACDC: The adverse conditions dataset with correspondences for semantic driving scene understanding. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)* (pp. 10765-10775). IEEE. <https://doi.org/10.1109/ICCV48922.2021.01061>
- [7] Sakaridis, C., Dai, D., & Van Gool, L. (2018). Semantic foggy scene understanding with synthetic data. *International Journal of Computer Vision*, 126(9), 973-992. <https://doi.org/10.1007/s11263-018-1072-8>
- [8] Hahner, M., Sakaridis, C., Dai, D., & Van Gool, L. (2021). Fog simulation on real LiDAR point clouds for 3D object detection in adverse weather. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)* (pp. 15283-15292). IEEE. <https://doi.org/10.1109/ICCV48922.2021.01489>
- [9] Hahner, M., Sakaridis, C., Bijelic, M., Heide, F., Yu, F., Dai, D., & Van Gool, L. (2022). LiDAR snowfall simulation for robust 3D object detection. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 16364-16374). IEEE. <https://doi.org/10.1109/CVPR52688.2022.01590>
- [10] Vora, S., Lang, A. H., Helou, B., & Beijbom, O. (2020). PointPainting: Sequential fusion for 3D object detection. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 4604-4612). IEEE. <https://doi.org/10.1109/CVPR42600.2020.00466>
- [11] Teng, D. (2025). PACO: Predictive auto-configuration for SLO-constrained large language model inference serving. *Innovation and Technology Studies*, 2(1), 1-10.
- [12] Yoo, J. H., Kim, Y., Kim, J., & Choi, J. W. (2020). 3D-CVF: Generating joint camera and LiDAR features using cross-view

- spatial feature fusion for 3D object detection. In Proceedings of the European Conference on Computer Vision (ECCV) (pp. 720-736). Springer. https://doi.org/10.1007/978-3-030-58589-8_43
- [13] Bai, X., Hu, Z., Zhu, X., Huang, Q., Chen, Y., Fu, H., & Tai, C.-L. (2022). TransFusion: Robust LiDAR-camera fusion for 3D object detection with transformers. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (pp. 1090-1099). IEEE. <https://doi.org/10.1109/CVPR52688.2022.00118>
- [14] Liu, Z., Tang, H., Amini, A., Yang, X., Mao, H., Rus, D., & Han, S. (2023). BEVFusion: Multi-task multi-sensor fusion with unified bird's-eye view representation. In Proceedings of the IEEE International Conference on Robotics and Automation (ICRA) (pp. 2774-2781). IEEE. <https://doi.org/10.1109/ICRA48891.2023.10160915>
- [15] Lang, A. H., Vora, S., Caesar, H., Zhou, L., Yang, J., & Beijbom, O. (2019). PointPillars: Fast encoders for object detection from point clouds. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (pp. 12697-12705). IEEE. <https://doi.org/10.1109/CVPR.2019.01298>
- [16] Teng, D. (2025). TEAS: Token- and energy-aware autoscaling for cost-efficient LLM serving. *AI and Data Science Journal*, 6(3), 1-10.
- [17] Zhou, Y., & Tuzel, O. (2018). VoxelNet: End-to-end learning for point cloud based 3D object detection. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (pp. 4490-4499). IEEE. <https://doi.org/10.1109/CVPR.2018.00472>
- [18] Qi, C. R., Yi, L., Su, H., & Guibas, L. J. (2017). PointNet++: Deep hierarchical feature learning on point sets in a metric space. In Advances in Neural Information Processing Systems 30 (NeurIPS) (pp. 5099-5108).
- [19] Yin, T., Zhou, X., & Krahenbuhl, P. (2021). Center-based 3D object detection and tracking. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (pp. 11784-11793). IEEE. <https://doi.org/10.1109/CVPR46437.2021.01161>
- [20] Tian, Z., Shen, C., Chen, H., & He, T. (2019). FCOS: Fully convolutional one-stage object detection. In Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) (pp. 9627-9636). IEEE. <https://doi.org/10.1109/ICCV.2019.00972>
- [21] Lin, T.-Y., Goyal, P., Girshick, R., He, K., & Dollár, P. (2017). Focal loss for dense object detection. In Proceedings of the IEEE International Conference on Computer Vision (ICCV) (pp. 2999-3007). IEEE. <https://doi.org/10.1109/ICCV.2017.324>
- [22] Kendall, A., & Gal, Y. (2017). What uncertainties do we need in Bayesian deep learning for computer vision? In Advances in Neural Information Processing Systems 30 (NeurIPS) (pp. 5574-5584).
- [23] Gal, Y., & Ghahramani, Z. (2016). Dropout as a Bayesian approximation: Representing model uncertainty in deep learning. In Proceedings of the 33rd International Conference on Machine Learning (ICML) (pp. 1050-1059). PMLR.
- [24] Zi, B. (2024). Large language models for enterprise workflow automation in financial operations. *Innovation and Technology Studies*, 1(1), 24-29.
- [25] Ding, J., Shen, Z., & Liu, W. (2026). Game-theoretic cost-sensitive adversarial training for robust cloud intrusion detection against GAN-based evasion attacks. *Applied Sciences*, 16(8), Article 3944. <https://doi.org/10.3390/app16083944>
- [26] Teng, D., Rhee, M., Qin, Y., Zi, B., & Liu, W. (2026). SW-SpeedDLM: Sliding-window speculative decoding for diffusion language models under long-context constraints. *Mathematics*, 14(12), Article 2137. <https://doi.org/10.3390/math14122137>
- [27] Zi, B. (2024). Cloud-native distributed systems for real-time payment intelligence. *AI and Data Science Journal*, 1(1), 51-56.
- [28] Chen, Z., Wang, M., Zeng, Z., & Ping, W. (2026). Uncertainty-aware financial forecasting: Leveraging conformal prediction for risk-adjusted models. *IEEE Access*, 14, 1-18.
- [29] Wang, B., Wang, Z., Zhao, W., Zhang, F., & Shang, W. (2026). DRL-Adapt: Deep reinforcement learning for adaptive routing convergence optimization in large-scale networks. *IEEE Open Journal of the Computer Society*, 7, 1-14.
- [30] Zhang, H. (2025). Physics-informed neural networks for high-fidelity electromagnetic field approximation in VLSI and RF EDA applications. *Journal of Computing and Electronic Information Management*, 18(2), 38-46.
- [31] Jiao, Y., Fan, H., Yue, X., Ping, W., Sun, T., & Wang, J. (2026). Dynamic heterogeneous graph contrastive learning for uncovering collusive financial fraud. *Scientific Reports*, 16, Article 14567. <https://doi.org/10.1038/s41598-026-58938-5>
- [32] Wang, Z., Yang, J. S., Shang, W., & Ding, J. (2026). FairPromote: Explainable and fairness-aware talent promotion prediction via adversarial debiasing and SHAP-based interpretation. *IEEE Access*, 14, 1-16.