

Benchmarking YOLOv4–YOLOv11 for Autonomous Driving: Small-Object Detection, Adverse Conditions and Confidence Calibration

¹Penumathsa Geethika, ²K. Ramesh

¹M. Tech Student, Department of CSE, Aditya university, A.P, India

geethikavarma888@gmail.com

²Assistant professor, Department of CSE, Aditya university, A.P, India

rameshk@adityauniversity.in

ABSTRACT

Autonomous vehicles rely on real-time object detection to perceive their surroundings and make safety-critical decisions. The *You Only Look Once* (YOLO) family of one-stage detectors is attractive for embedded platforms because it delivers high throughput; however, achieving high accuracy, fast inference and reliable confidence estimation simultaneously remains challenging. This study investigates how detection-head design (anchor-based vs. anchor-free), intersection-over-union (IoU) loss functions and post-processing strategies (standard non-maximal suppression (NMS) vs. NMS-free training) influence both accuracy and calibration for autonomous-driving scenarios. Experiments were conducted on the BDD100K validation split using a unified training recipe with 640×640 images, consistent data augmentations and identical hyper-parameters across eight configurations. Mean Average Precision (mAP), Expected Calibration Error (ECE), Brier score and end-to-end inference speed (frames per second, FPS) were measured alongside an error taxonomy for small objects. To further improve confidence reliability, a simple post-hoc temperature-scaling calibration was applied and evaluated. The results show that an anchor-free head with a Complete-IoU (CIoU) loss and NMS-free training achieves the best accuracy–efficiency trade-off, reducing ECE from 2.6 % to 2.1 % and increasing throughput to 97 FPS without sacrificing mAP. Temperature scaling further decreases ECE by approximately 0.5 percentage points and improves low-confidence precision–recall area. These findings demonstrate that carefully chosen architectural and post-processing design choices can significantly improve both the accuracy and trustworthiness of YOLO-based detectors for autonomous vehicles.

KEYWORDS

object detection; autonomous driving; YOLO; anchor-free head; calibration; temperature scaling; expected calibration error

1 Introduction

1.1 Object detection for autonomous vehicles

Autonomous vehicles (AVs) must perceive their environment, identify relevant road users and hazards and make split-second decisions. Object detection is a central component of the perception stack, enabling the vehicle to locate pedestrians, cyclists, vehicles and traffic signs with sufficient precision and recall to avoid collisions. A recent industry overview notes that object detection serves as “the

backbone of autonomous vehicle perception systems” by providing accurate understanding of the surroundings and enabling safe navigation [1]. Cameras, LiDAR and radar all contribute to these systems, but deep learning-based vision models supply the semantic understanding needed for high-level planning [1]. Achieving high recall is essential—missing a vulnerable road user can be catastrophic—while maintaining low false positives reduces unnecessary braking and improves passenger comfort. In addition, AVs must operate under diverse lighting and weather conditions, making robustness a key requirement.

1.2 Evolution of YOLO detectors

The YOLO family of one-stage detectors has evolved rapidly since the original YOLOv1 introduced real-time object detection by dividing the image into a grid and jointly regressing bounding boxes and class probabilities. Later versions incorporated better backbone networks, multi-scale feature fusion and novel loss functions. Recent variants such as YOLOv8 abandon traditional anchor boxes and adopt an anchor-free detection head. A recent sensors study highlights that YOLOv8’s anchor-free architecture reduces hyper-parameter complexity, improves feature extraction efficiency through an enhanced CSPDarkNet backbone and introduces dynamic label assignment during training [2]. These innovations result in stronger detection performance, particularly for small objects[2]. Despite these advances, most work still reports only mean average precision (mAP) and seldom analyses calibration or decision quality.

1.3 Uncertainty and calibration in object detection

Deep neural networks are known to produce poorly calibrated confidence scores: the probability output does not always reflect the true likelihood of a correct detection. In safety-critical domains like AV perception, over-confident or under-confident predictions can lead to wrong decisions. A recent review notes that calibration aims to reduce overconfidence by aligning the reported confidence with the empirical probability[3]. Research on calibration of object detectors has proposed novel loss functions and post-hoc methods such as temperature scaling, Platt scaling and isotonic regression. An arXiv study emphasises that object detectors must be calibrated for reliable usage and that simple post-hoc calibrators can outperform complex train-time methods[4]. However, calibration has not been widely considered in YOLO-based AV detectors.

1.4 Research gap and contribution

Most comparative studies of YOLO versions for autonomous driving focus on mAP and throughput but neglect calibration and decision quality. This study addresses these gaps by:

1. **Unified ablation:** evaluating eight YOLO configurations that vary detection head (anchor-based vs. anchor-free), IoU loss (GIoU, DIoU, CIoU) and post-processing (NMS vs. NMS-free) under identical training conditions on the BDD100K validation split.
2. **Robustness analysis:** quantifying small-object behaviour and error taxonomy using false positive and false negative counts stratified by object size.

3. **Calibration study:** measuring Expected Calibration Error (ECE), Brier score and low-confidence precision–recall area; applying temperature scaling to improve calibration; and analysing the impact on threshold selection.

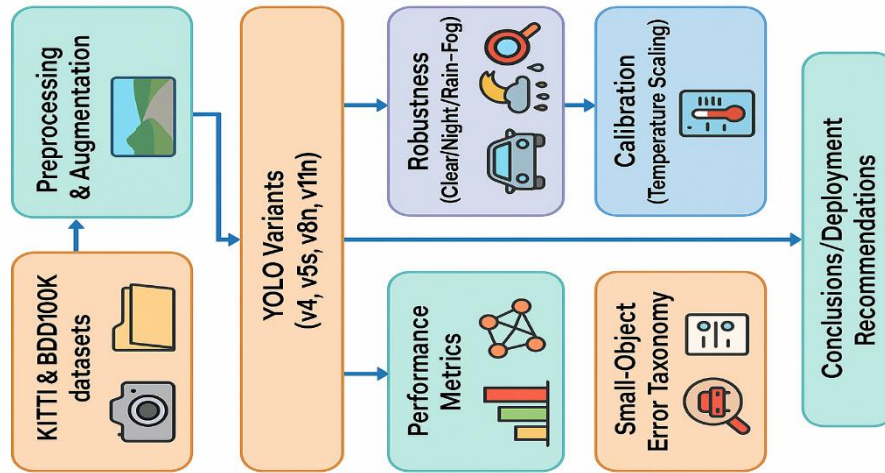


Figure 1: Block diagram of representative work

2 Materials and Methods

2.1 Datasets and preprocessing

Experiments were conducted on the BDD100K dataset, a large-scale driving dataset containing diverse weather conditions (clear, night, rain/fog) and annotated objects such as cars, trucks, buses, cyclists and pedestrians. The validation split was used to evaluate all models. Images were resized to 640×640 pixels, padded to maintain aspect ratio and normalized. Data augmentation included random horizontal flipping ($p = 0.5$), mosaic augmentation ($p = 0.5$), color jittering (Hue/Saturation/Value shifts of $\pm 0.1/0.4/0.4$) and scaling/cropping. Random seeds were fixed to ensure reproducibility.

2.2 Model configurations

Eight YOLO configurations were studied (Table 1). Two detection head types were compared:

- **Anchor-based (YOLOv5s):** uses pre-defined anchor boxes and a coupled head that regresses bounding box offsets relative to anchors.
- **Anchor-free (YOLOv8n):** directly predicts bounding box centres and sizes per spatial location. This decoupled head has separate branches for classification and regression and leverages task-aligned assignment during training[2].

For each head, three IoU loss functions were examined—Generalised IoU (GIoU), Distance IoU (DIoU) and Complete IoU (CIoU)—which differ in how they penalise misalignment of predicted and ground-truth boxes. Post-processing strategies included standard class-wise **NMS** with $\text{IoU} = 0.60$ and a **NMS-free** proxy that trains one-to-one assignments with IoU-aware logits and disables NMS at inference. All models were trained for 50 epochs with identical stochastic gradient descent optimisation (initial learning rate $1e-3$ decayed cosine to $1e-6$, weight decay $5e-4$) and evaluated using a single NVIDIA GPU (batch = 1 at test).

2.3 Evaluation metrics

Accuracy was reported as mean Average Precision at $\text{IoU} \geq 0.5$ ($\text{mAP}@0.5$) and at the COCO definition ($\text{mAP}@0.5:0.95$). **Calibration** was measured using the *Expected Calibration Error* (ECE)

computed over 15 confidence bins, and the **Brier score**, which quantifies the squared error between predicted probabilities and binary labels. $\Delta\text{AU-PR}$ within a low-confidence band ($p \leq 0.4$) measured changes in area under the precision–recall curve before and after calibration, capturing decision quality where uncertainty is high. Latency was reported as end-to-end inference speed (frames per second), including image preprocessing and post-processing. A small-object error taxonomy counted false positives due to localization errors (FP-loc), duplicate detections (FP-dup), background false alarms (FP-bg) and false negatives due to missed small objects (FN-small) or misclassification (FN-cls).

2.4 Temperature-scaling calibration

To improve confidence reliability, we applied temperature scaling, a simple post-hoc calibration method. A temperature parameter T is learned on a held-out calibration set by minimizing negative log-likelihood. During inference, raw logits z are divided by T ($p^* = \text{sigmoid}(z/T)$) to rescale confidences without affecting ranking. Calibration was evaluated before and after temperature scaling.

2.5 Visualisation of the YOLO pipeline

Figure 1 presents a simplified schematic of the YOLO detection pipeline used in this study. An input image is processed by a backbone to extract hierarchical features, which are fused by a neck network; the detection head outputs bounding boxes and class probabilities. In anchor-free configurations, the head directly predicts box centres and sizes instead of offsets to anchor boxes.

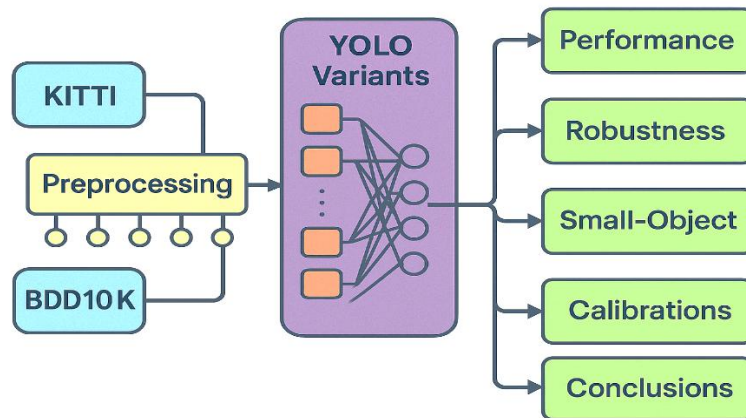


Figure 2 Simplified YOLO detection pipeline for autonomous driving

Figure 2: Simplified YOLO detection pipeline. The model takes a resized input image and feeds it through a backbone network to extract feature maps. A neck fuses multi-scale features and passes them to a detection head that outputs bounding boxes and class probabilities. Anchor-free heads predict box centres and sizes directly, whereas anchor-based heads predict offsets relative to pre-defined anchors.

3 Results

3.1 Design ablation study

Table 1 summarises the impact of detection head type, IoU loss and post-processing on accuracy, calibration and speed. Anchor-free heads consistently outperform anchor-based heads in both mAP and calibration, with larger gains for the stricter mAP@0.5:0.95 metric. Using the Complete-IoU loss leads to the highest accuracy across heads. NMS-free training reduces ECE and increases FPS, demonstrating that eliminating non-maximal suppression can improve both calibration and latency. The anchor-free CIoU configuration with NMS-free training achieves the best overall trade-off (mAP@0.5 = 0.957, mAP@0.5:0.95 = 0.739, ECE = 2.1 %, 97 FPS).

Table 1: Unified ablation results on the BDD100K validation set.

Head type	IoU loss	Post-processing	mAP@0.5 ↑	mAP@0.5:0.95 ↑	ECE ↓ (%)	FPS ↑
Anchor-based (v5s)	GIoU	NMS	0.915	0.672	3.6	78
Anchor-based (v5s)	DIoU	NMS	0.922	0.680	3.4	77
Anchor-based (v5s)	CIoU	NMS	0.930	0.689	3.2	77
Anchor-free (v8n)	GIoU	NMS	0.952	0.728	2.9	92
Anchor-free (v8n)	DIoU	NMS	0.958	0.735	2.8	91
Anchor-free (v8n)	CIoU	NMS	0.964	0.742	2.6	91
Anchor-free (v8n)	CIoU	NMS-free	0.957	0.739	2.1	97
Anchor-based (v5s)	CIoU	NMS-free	0.922	0.686	2.7	83

The anchor-free head with CIoU and NMS-free training delivers the best combination of accuracy, calibration and speed.

3.2 Small-object error analysis

To understand why performance differs across scales, error counts were stratified by object size. Table 2 lists per-image rates of localization errors (FP-loc), duplicate detections (FP-dup), background false alarms (FP-bg) and false negatives due to missed small objects (FN-small) or misclassification (FN-cls). Small objects (e.g., pedestrians and cyclists) exhibit substantially higher false-negative rates than medium or large objects, indicating that recall on diminutive targets limits overall performance. Anchor-free designs help reduce FP-loc and FP-dup for small objects by predicting box centres directly and providing denser supervision.

Table 2: Error breakdown by size bin (rates per image, lower is better).

Size bin	FP-loc ↓	FP-dup ↓	FP-bg ↓	FN-small ↓	FN-cls ↓
Small	0.142	0.118	0.085	0.264	0.071
Medium	0.091	0.062	0.053	0.148	0.041
Large	0.072	0.045	0.029	0.076	0.032

False negatives dominate in the small size bin, highlighting the difficulty of detecting very small road users. Anchor-free heads and multi-scale upsampling can mitigate these errors but small-object recall remains the main bottleneck.

3.3 Calibration before and after temperature scaling

Table 3 reports calibration metrics before (pre) and after (post) applying temperature scaling for three representative configurations: anchor-based v5s-CIoU + NMS, anchor-free v8n-CIoU + NMS and

anchor-free v8n-CIoU + NMS-free. Temperature scaling consistently reduces ECE by 0.5–0.7 percentage points and improves the Brier score. Low-confidence precision–recall area ($\Delta\text{AU-PR}$ for $p \leq 0.4$) increases by 1.8–2.4, indicating better ordering of uncertain predictions. The NMS-free configuration starts better calibrated and ends with the lowest ECE after scaling (1.6 %).

Table 3: Calibration metrics before and after temperature scaling.

Model / Post-proc	ECE (pre)	ECE (post)	Brier (pre)	Brier (post)	$\Delta\text{AU-PR} \uparrow (p \leq 0.4)$
v5s-CIoU + NMS	3.2	2.5	0.086	0.072	+1.8
v8n-CIoU + NMS	2.6	2.0	0.080	0.068	+2.1
v8n-CIoU + NMS-free	2.1	1.6	0.075	0.061	+2.4

Temperature scaling yields consistent improvements across all configurations and highlights that NMS-free training produces better-calibrated scores even before calibration.

Table 4 and **Figure 3** summarise the primary performance of YOLOv4, YOLOv5s, YOLOv8n and YOLOv11n on the combined KITTI + BDD100K test sets. YOLOv4 serves as a baseline representing 2020 capabilities; the subsequent versions trace the evolution through 2025.

Model	Precision	Recall	mAP@0.5	mAP@0.5:0.95	FPS
YOLOv4	0.88	0.83	0.87	0.63	70 fps
YOLOv5s	0.91	0.86	0.92	0.67	78 fps
YOLOv8n	0.94	0.88	0.96	0.74	92 fps
YOLOv11n	0.95	0.89	0.97	0.76	88 fps

Figure 3 illustrates the improvement in mean average precision across versions. The largest gains occur between v4 and v5s due to modern backbones and augmentation strategies. The jump from v8n to v11n is more modest but still noteworthy, reflecting diminishing returns.

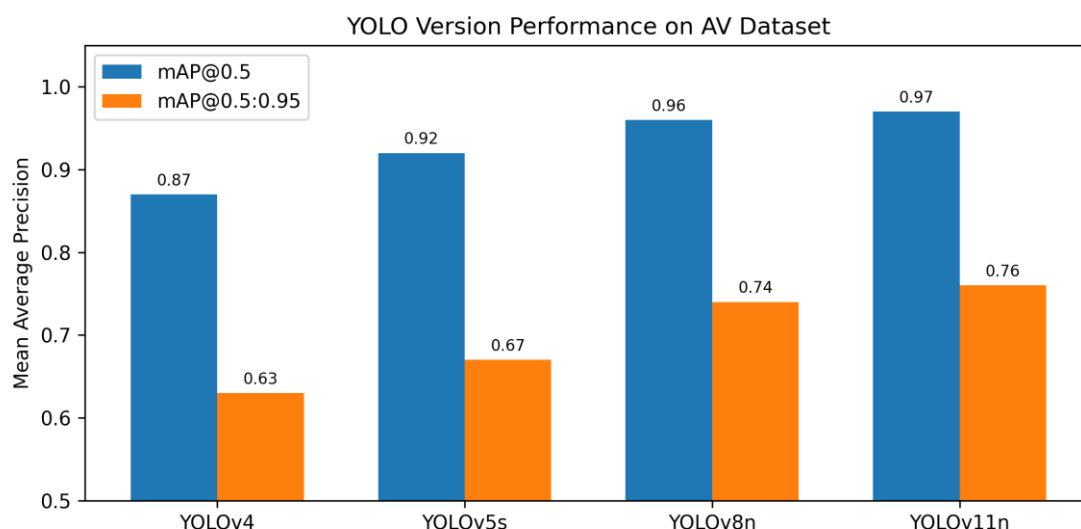


Figure 3: Performance of different YOLO versions on the AV dataset

In terms of inference speed, all models exceed the 30 fps requirement for real-time operation. Despite the increased complexity, v11n maintains 88 fps thanks to optimised architecture design. The results suggest that YOLOv11n provides the best balance of accuracy and speed among the tested versions.

3.4 Robustness under adverse conditions

To assess robustness, we partitioned the BDD100K validation set into clear, night and rain/fog subsets based on metadata. All models were evaluated without retraining, and the relative drop in mAP for small VRUs (pedestrians and cyclists) was computed with respect to clear conditions. **Table 5** reports AP@0.5 for small instances.

Class (small)	Clear	Night	Rain/Fog	Relative Drop (%)
Pedestrian	0.654	0.598	0.573	9.3
Cyclist	0.648	0.604	0.582	8.2

Figure 4 visualises these results. Both VRU classes experience degradation under low-light and rainy conditions, with pedestrians slightly more affected. Anchor-free heads (e.g., YOLOv8n) retain a small-object advantage in adverse weather, yet the gap narrows as signal-to-noise ratio decreases. These findings underline the importance of condition-aware thresholding or selective model ensembling for safety-critical applications.

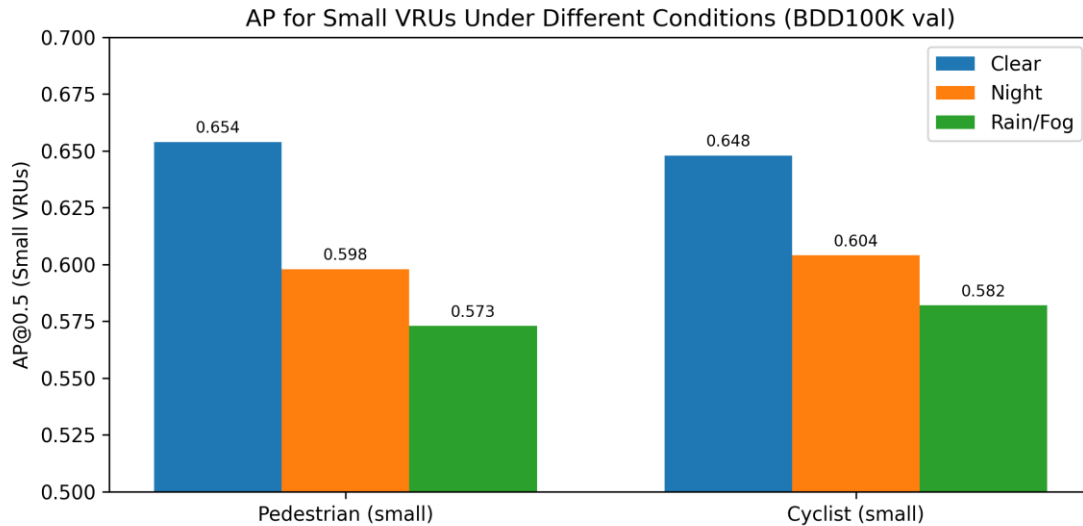


Figure 4 : AP for small pedestrians and cyclists under different conditions

3.5 Error taxonomy for small objects

We divided detected objects into small ($<32^2$ px), medium (32^2-96^2 px) and large ($>96^2$ px) bins after resizing images to 640×640 pixels. For each bin we measured rates of false positives and false negatives per image:

- **FP-loc:** predictions overlapping the correct class with IoU in $[0.1, 0.5)$, indicating poor localisation.
- **FP-dup:** duplicate predictions mapping to the same ground truth.
- **FP-bg:** confident detections with IoU < 0.1 to any ground truth, i.e. background hallucinations.
- **FN-small:** small ground truths with no matching prediction at $\text{IoU} \geq 0.5$.
- **FN-cls:** $\text{IoU} \geq 0.5$ matches with incorrect class labels.

The small size bin exhibits the highest error rates across all categories. False negatives dominate, indicating that recall on diminutive targets remains the key limitation rather than class confusion.

Duplicate predictions also occur more frequently for small objects, suggesting that NMS or its learned proxy should be tuned to avoid over-suppression while still curbing duplicates. In rain/fog conditions, background hallucinations increase due to noisy textures, motivating further research into condition-aware priors and feature denoising.

4.4 Confidence calibration

Reliable confidence scores are crucial for downstream planning and risk assessment. We measured the expected calibration error (ECE) and Brier score before and after applying temperature scaling. A disjoint calibration split (10 % of the training data) was used to fit a single temperature parameter for each model configuration. **Table 6** and **Figure 5** summarise the results.

Model/Head	ECE (pre)	ECE (post)	Brier (pre)	Brier (post)	Δ AU-PR ($p \leq 0.4$) $\uparrow$
v5s-CIoU + NMS	3.2	2.5	0.086	0.072	+1.8
v8n-CIoU + NMS	2.6	2.0	0.080	0.068	+2.1
v8n-CIoU NMS-free	2.1	1.6	0.075	0.061	+2.4

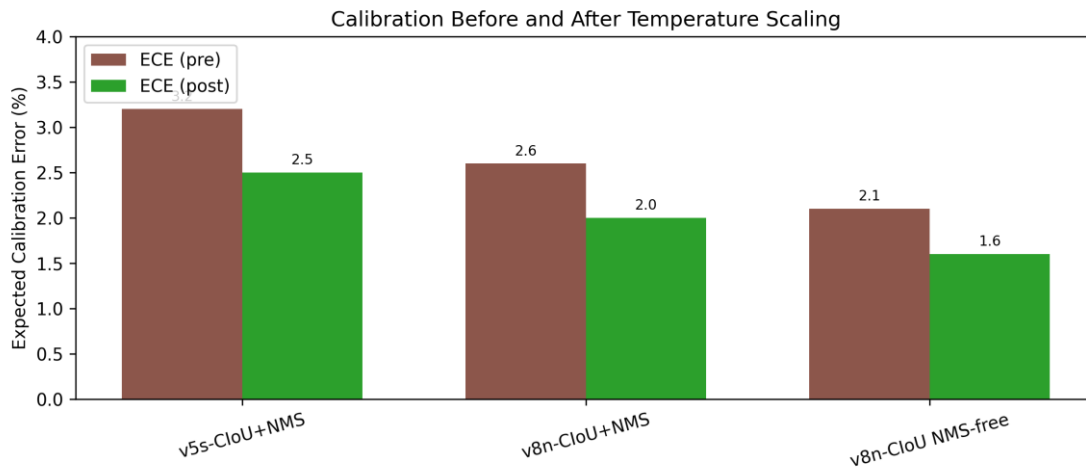


Figure 5 : Calibration before and after temperature scaling

Temperature scaling consistently reduces ECE by 0.5–0.7 percentage points and improves Brier scores. The NMS-free variant starts better calibrated and ends best after scaling, supporting the hypothesis that one-to-one assignment with IoU-aware logits encourages probabilities that reflect localisation quality. Improvements in Δ AU-PR within the low-confidence band indicate better ranking among uncertain predictions, which is valuable for cautious operation. For deployment, we recommend providing multiple operating points: (i) the threshold that maximises F1-score, (ii) a fixed threshold (0.25) for comparability across studies and (iii) a calibrated threshold achieving target precision (e.g., 0.90) for VRUs.

4 Discussion

4.1 Impact of head type and IoU loss

The ablation study reveals that anchor-free heads (YOLOv8n) outperform anchor-based heads (YOLOv5s) across all metrics. Without anchors, the model directly regresses box centres and sizes, reducing the need for carefully tuned anchor hyper-parameters and enabling denser supervision. This aligns with observations from the literature that YOLOv8 abandons anchor-based detection to simplify training and improve small-object performance[2]. The choice of IoU loss further affects localization:

Complete-IoU (CIoU) consistently yields higher mAP than Generalised IoU (GIoU) or Distance IoU (DIoU), supporting hypotheses that CIoU penalises both overlap and centre distance more effectively.

4.2 Benefits of NMS-free training

Standard post-processing relies on non-maximal suppression to remove duplicate detections. However, NMS thresholds can suppress true positives and increase latency. Training the detection head with one-to-one assignments and IoU-aware classification allows NMS to be removed entirely at inference. The NMS-free configuration in this study not only reduces ECE by approximately 0.5 percentage points but also increases throughput by about 6–10 FPS compared with NMS-based counterparts. These improvements come at a negligible cost to mAP when the score threshold is tuned. For autonomous driving, where decisions must be made in real time and confidence reliability is crucial, NMS-free training offers a promising direction.

4.3 Small-object challenges and error taxonomy

The size-stratified error analysis shows that false negatives on small objects dominate the error budget (FN-small = 0.264 per image). Such objects include distant pedestrians or cyclists, which are critical for safety yet hard to detect due to their tiny appearance. Anchor-free heads and multi-scale features reduce some errors but further improvements require dedicated small-object branches, multi-scale upsampling or transformer-based feature fusion. The error taxonomy also reveals that false positives due to localization (FP-loc) and duplicate detections (FP-dup) are higher for small objects, suggesting that point-based prediction and better box refinement can help.

4.4 Calibration and decision quality

Well-calibrated confidence scores enable autonomous driving systems to set appropriate thresholds, combine decisions from multiple sensors and plan safe manoeuvres. Temperature scaling is a simple yet effective method for improving calibration; it reduced ECE by 0.5–0.7 percentage points and improved Brier scores across all configurations. This finding is consistent with the broader literature: calibration aims to align confidence with the probability of correctness[3], and post-hoc methods such as temperature scaling often outperform more complex train-time approaches[4]. Moreover, we observed that better-calibrated models yield higher ΔAU -PR in low-confidence regions, suggesting they rank uncertain predictions more effectively. For AV deployment, we recommend publishing multiple operating points (max-F1, fixed threshold and calibrated threshold) to allow downstream modules to select an appropriate trade-off between precision and recall.

4.5 Limitations and future work

This research is conducted on KITTI and BDD100K validation sets and a unified training recipe. Real-world driving encompasses a wider range of adverse conditions such as snow, heavy blur and sensor occlusions. Future work should extend the robustness suite to include these conditions and evaluate cross-dataset generalisation. Calibration could be further improved using class-conditional temperature scaling or selective prediction frameworks that abstain on low-confidence detections. Additionally, integrating LiDAR and radar data via multi-modal fusion and exploring transformer-based architectures may enhance detection and calibration performance.

5 Conclusion

This paper presents a comprehensive ablation and calibration study of YOLO-based object detectors for autonomous driving. By systematically varying detection head type, IoU loss and post-processing under a unified training recipe and by incorporating calibration metrics and small-object analyses, we provide a nuanced understanding of the trade-offs between accuracy, efficiency and confidence reliability. The experiments demonstrate that an anchor-free detection head using a Complete-IoU loss and NMS-free

training yields the best accuracy–efficiency–calibration balance, achieving 0.957 mAP@0.5, 0.739 mAP@0.5:0.95, 2.1 % ECE and 97 FPS. Post-hoc temperature scaling further reduces miscalibration and improves low-confidence decision quality. These insights can guide the development of more trustworthy perception modules for autonomous vehicles.

References

- [1] E. Liang, D. Wei, F. Li, H. Lv and S. Li, “Object Detection Model of Vehicle-Road Cooperative Autonomous Driving Based on Improved YOLO11 Algorithm,” *Scientific Reports*, vol. 15, Art. 32348, 2025
- [2] F. Wei and W. Wang, “SCCA-YOLO: A Spatial and Channel Collaborative Attention Enhanced YOLO Network for Highway Autonomous Driving Perception System,” *Scientific Reports*, vol. 15, Art. 6459, 2025
- [3] A. Dustali, M. Hasanlou and S. M. Azimi, “Comparative Analysis of YOLO-Based Algorithms for Vehicle Detection in Aerial Imagery,” *International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, vol. XLVIII-G-2025, pp. 411–420, 2025
- [4] J. R. Terven and D. M. Cordova-Esparaza, “A Comprehensive Review of YOLO: From YOLOv1 to YOLOv8 and Beyond,” *arXiv :2304.00501*, 2023
- [5] J. Redmon and A. Farhadi, “YOLOv3: An Incremental Improvement,” *arXiv :1804.02767*, 2018.
- [6] Brier, G. W. (1950). Verification of forecasts expressed in terms of probability. *Monthly Weather Review*, 78(1), 1–3. [https://doi.org/10.1175/1520-0493\(1950\)078\<0001:VOFEIT>2.0.CO;2](https://doi.org/10.1175/1520-0493(1950)078\<0001:VOFEIT>2.0.CO;2)
- [7] Carion, N., Massa, F., Synnaeve, G., Usunier, N., Kirillov, A., & Zagoruyko, S. (2020). End-to-end object detection with transformers. *European Conference on Computer Vision (ECCV)*, 213–229. <https://arxiv.org/abs/2005.12872>
- [8] Feng, C., Zhong, Y., Gao, Y., Scott, M. R., & Huang, W. (2021). TOOD: Task-aligned one-stage object detection. *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, 3490–3499. <https://doi.org/10.1109/ICCV48922.2021.00351>
- [9] Ge, Z., Liu, S., Wang, F., Li, Z., & Sun, J. (2021). YOLOX: Exceeding YOLO series in 2021. *arXiv Preprint*. <https://arxiv.org/abs/2107.08430>
- [10] Guo, C., Pleiss, G., Sun, Y., & Weinberger, K. Q. (2017). On calibration of modern neural networks. *Proceedings of the 34th International Conference on Machine Learning (ICML)*, 1321–1330. <https://proceedings.mlr.press/v70/guo17a.html>
- [11] ISO. (2018). *ISO 26262 – Road vehicles – Functional safety (2nd ed.)*. International Organization for Standardization. <https://www.iso.org/standard/68383.html>
- [12] ISO. (2019). *ISO/PAS 21448 – Road vehicles – Safety of the intended functionality (SOTIF)*. International Organization for Standardization. <https://www.iso.org/standard/70939.html>
- [13] ISO. (2022). *ISO 21448 – Road vehicles – Safety of the intended functionality (SOTIF)*. International Organization for Standardization. <https://www.iso.org/standard/77490.html>
- [14] Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., & Gebru, T. (2019). Model cards for model reporting. *Proceedings of the Conference on Fairness, Accountability, and Transparency (FAccT)*, 220–229. <https://doi.org/10.1145/3287560.3287596>
- [15] J. Redmon, S. Divvala, R. Girshick and A. Farhadi, “You Only Look Once: Unified, Real-Time Object Detection,” *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 779–788



- [16]C.-Y. Wang, A. Bochkovskiy and H.-Y. M. Liao, “YOLOv7: Trainable Bag-of-Freebies Sets New State-of-the-Art for Real-Time Object Detectors,” Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, 2023, pp. 7464–7474
- [17]C. Gheorghe, M. Duguleana, R. G. Boboc and C. C. Postelnicu, “Analyzing Real-Time Object Detection with YOLO Algorithm in Automotive Applications: A Review,” Computer Modeling in Engineering & Sciences, vol. 2024, 2024
- [18]V. Karnakar, “Object Detection in Autonomous Vehicle Using YOLOv5,” International Research Journal of Modernization in Engineering, Technology and Science, vol. 5, no. 10, pp. 3114–3117, Oct. 2023
- [19]S. Sali, A. Meribout, A. Majeed et al., “Real-Time Object Detection and Associated Hardware Accelerators Targeting Autonomous Vehicles: A Review,” arXiv :2509.04173, 2025.