

# *An Efficient Visual Lane Detection Method for Complex Road Scenes*

Caixia Zhao<sup>1</sup>, Liang Shao<sup>2</sup>

<sup>1</sup>*School of Air Transportation and Engineering, Nanhang Jincheng College, Nanjing, China*

<sup>2</sup>*AVIC General Research Institute Co.,Ltd, Zhuhai, China*

**Keywords:** Lane Detection; Deep Learning; Context Feature; Geometric Consistency

**Abstract:** Lane detection is essential for autonomous driving and advanced driver-assistance systems, yet remains challenging in complex road scenes with occlusion, illumination variations, adverse weather, and ambiguous lane markings. To address these issues, we propose **Context Feature Guided CLRNet (CFG-CLRNet)**, a vision-based lane detection framework built upon CLRNet. Specifically, **Context Feature Guidance Module (CFGM)** is introduced to enhance discriminative lane features by integrating coordinate attention with dual-branch atrous convolution, enabling effective multi-scale contextual interaction while preserving position-sensitive information. In addition, **Geometric Consistency Loss (GCL)** is designed to regularize lane geometry through second-order differences of lane points, improving structural continuity and geometric localization under incomplete or degraded visual observations. Extensive experiments on CULane and TuSimple demonstrate that CFG-CLRNet consistently improves lane detection performance over the baseline and exhibits stronger robustness in challenging road conditions, including occlusion, nighttime scenes, and curved lanes.

## 1. Introduction

Lane detection plays a central role in autonomous driving and advanced driver-assistance systems by providing structured road geometry for localization, motion planning, and vehicle control. Recent learning-based methods have largely replaced conventional pipelines built on handcrafted features and geometric assumptions, leading to substantial gains in detection accuracy and runtime efficiency. However, robust lane perception remains difficult in unconstrained road environments, where lane markings may be incomplete, visually ambiguous, or severely corrupted by occlusion, illumination changes, shadows, glare, and adverse weather.

Existing deep lane detectors can be broadly categorized into segmentation-based, curve-based, and anchor-based methods. Segmentation-based approaches [1,2] formulate lane detection as dense pixel-wise classification and are effective at preserving fine-grained lane structures. However, they typically require high-resolution feature maps and additional instance grouping, resulting in considerable computational overhead. Curve-based methods [3,4] employ compact parametric representations to reduce prediction redundancy, but limited curve models may struggle with irregular geometries and complex lane topologies. Anchor-based methods [5] represent lanes using predefined or learnable anchors and regress their geometric offsets. Although efficient, their

performance depends heavily on anchor coverage and proposal quality, particularly for highly curved, densely distributed, or partially visible lanes.

Recent approaches, such as CLRNNet [6] and its extension [13], improve lane detection through hierarchical feature refinement, proposal-aligned feature aggregation, and lane-aware optimization. However, adverse conditions can substantially weaken local lane evidence, while recovering complete lane structures requires effective contextual and geometric reasoning. Increasing feature resolution or introducing heavy global interaction modules may improve robustness but often compromises inference efficiency. Therefore, balancing fine-grained localization, structural consistency, and computational efficiency remains a central challenge for practical lane detection.

To address the above challenges, this paper proposes **Context Feature Guided CLRNNet (CFG-CLRNNet)**, a context-enhanced lane detection framework built upon CLRNNet. The proposed method aims to jointly improve contextual feature modeling and lane geometric representation. Specifically, we design **Context Feature Guidance Module (CFGM)** in the feature extraction stage, which integrates coordinate attention with dual-branch atrous convolution. By preserving position-sensitive information while enlarging the effective receptive field, CFGM promotes multi-scale contextual interaction and enhances the discriminative representation of slender lane structures under complex backgrounds. In addition, we introduce **Geometric Consistency Loss (GCL)** during network optimization. GCL imposes a geometric smoothness constraint based on the second-order differences of lane points, encouraging the model to learn lane curves that better conform to realistic road geometry and alleviating discontinuous predictions caused by occlusion or local missing observations.

Extensive experiments are conducted on two public benchmarks, CULane and TuSimple, to evaluate the effectiveness of the proposed method. Experimental results show that CFG-CLRNNet consistently outperforms the baseline CLRNNet in detection accuracy and robustness. In particular, CFG-CLRNNet achieves more stable performance in challenging scenarios such as occlusion, nighttime scenes, and curved lanes, demonstrating its effectiveness and generalization capability in complex road environments.

## 2. Related Work

Existing deep learning-based lane detection methods can be broadly categorized into segmentation-based, curve-based, and anchor-based approaches according to their lane representations and prediction strategies.

**Segmentation-based methods** formulate lane detection as a dense pixel-wise classification task that separates lane regions from the background. This formulation provides flexible representations without imposing explicit geometric assumptions on lane shapes. SCNN [7] introduces a spatial message-passing mechanism within a semantic segmentation framework, allowing information to propagate across rows and columns to capture long-range dependencies and improve the continuity of partially visible lanes. RESA [2] further strengthens contextual modeling through recurrent feature shifting and aggregation, enabling more efficient information exchange over the feature map. Other segmentation-based methods commonly incorporate multi-scale features or instance-aware representations to distinguish adjacent lane markings. Although these approaches provide fine-grained localization and can accommodate lanes with diverse geometries, their reliance on dense high-resolution prediction typically incurs considerable computational cost. Moreover, additional clustering or post-processing may be required to separate individual lane instances, which can introduce errors in crowded scenes or under severe occlusion.

**Curve-based methods** describe each lane using a compact set of curve parameters and directly regress the corresponding geometric representation. LSTR [8] employs a Transformer architecture

to capture global image context and model the long-range structure of lane markings before predicting curve parameters. PolyLaneNet [4] represents lanes using polynomial functions, while other methods [9] adopt Bézier curves to provide a continuous and naturally smooth lane representation. Compared with dense segmentation, curve-based methods substantially reduce output redundancy and generally require little or no instance-grouping post-processing. Their compact representations also facilitate end-to-end lane prediction. Nevertheless, the localization of an entire lane depends on a small number of regressed parameters, making the prediction sensitive to parameter inaccuracies. In addition, predefined low-dimensional curve models may have insufficient flexibility to represent abrupt curvature changes, lane branches, intersections, or other complex road topologies.

**Anchor-based methods** detect lanes relative to predefined or dynamically generated line and row anchors. Line-CNN [5] formulates lane detection as line proposal generation followed by proposal refinement, directly producing structured lane instances rather than dense segmentation maps. LaneATT [11] extracts anchor-aligned features and models interactions among lane proposals to incorporate global contextual information. SGNet [14] introduces geometric guidance from the vanishing point to improve the consistency of lane localization. Row-anchor-based methods discretize lane positions at a predefined set of image rows. UFLD [12] reformulates lane localization as row-wise classification to achieve a simple and efficient detection pipeline, whereas CondLaneNet [10] employs conditional convolution to generate instance-specific lane predictions and better distinguish closely distributed lanes. Despite their effectiveness, fixed anchor distributions and sampling patterns may provide insufficient coverage for highly curved, nearly horizontal, partially visible, or densely distributed lanes. CLRNet [6] addresses this issue through cross-layer refinement, where high-level features provide semantic and contextual cues for initial lane detection, while low-level features progressively refine lane coordinates. Its subsequent extension [13] further improves proposal feature aggregation and the discrimination of complex lane structures. Anchor-based methods have therefore become a competitive paradigm, although their performance remains dependent on anchor initialization, proposal assignment, and feature sampling quality.

### 3. Methodology

#### 3.1. Overall Architecture

The proposed CFG-CLRNet adopts a top-down pipeline comprising a backbone network, a Feature Pyramid Network (FPN), a Context Feature Guided Module (CFGM), a cross-layer refinement head (ROIGather), and a multi-task optimization objective. The overall architecture is illustrated in Fig. 1.

Given an input image  $I \in R^{H \times W \times 3}$ , it is first resized to  $320 \times 800$  and fed into a backbone network (DLA34 or ResNet) to extract multi-scale features  $\{C_3, C_4, C_5\}$  with downsampling rates of  $\{8, 16, 32\}$ , respectively. A standard top-down FPN is subsequently applied to produce fused feature pyramids  $\{P_3, P_4, P_5\}$ , each with 64 channels.

The core of the proposed method lies in the Context Feature Guided Module (CFGM), which is inserted after each FPN level. CFGM first applies Coordinate Attention (CA) to recalibrate channel responses along horizontal and vertical directions, followed by dual-branch dilated convolutions with dilation rates  $r = 1$  and  $r = 3$  to capture multi-scale contextual information. The outputs are concatenated along the channel dimension, fused by a  $1 \times 1$  convolution, and added back via a residual connection. After CFGM enhancement, the refined features are fed into the ROIGather module for cross-layer refinement, followed by the detection head for classification and regression.

During training, the SimOTA dynamic label assignment strategy is adopted, and the network is optimized with a multi-task loss comprising the Focal loss, Line IoU loss, and the proposed Geometry Consistency Loss (GCL).

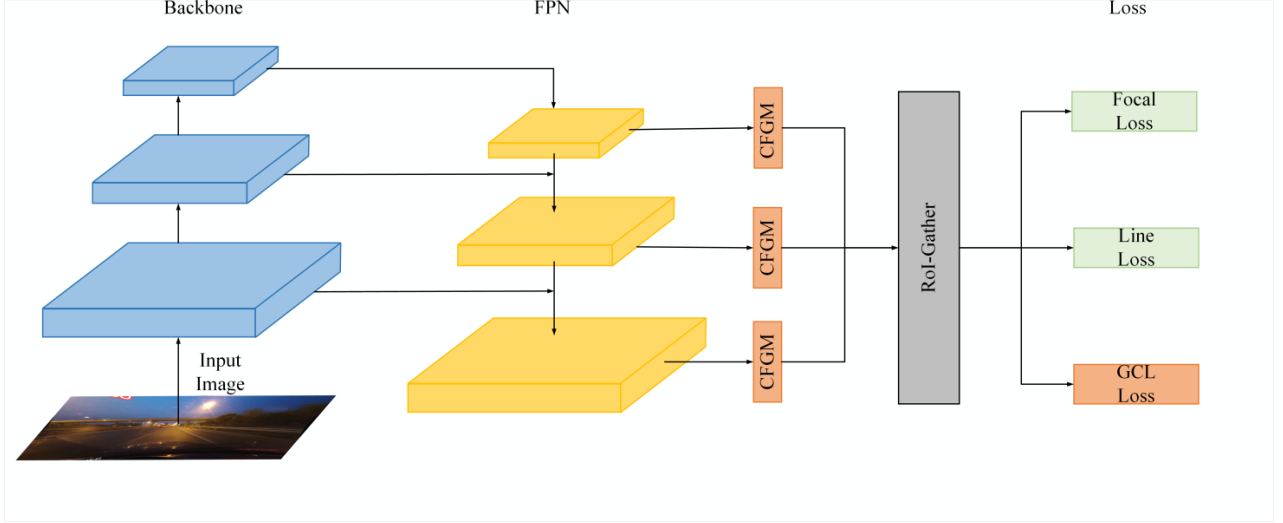


Figure 1: Architecture of CFG-CLRNet

### 3.2. Context Feature Guided Module

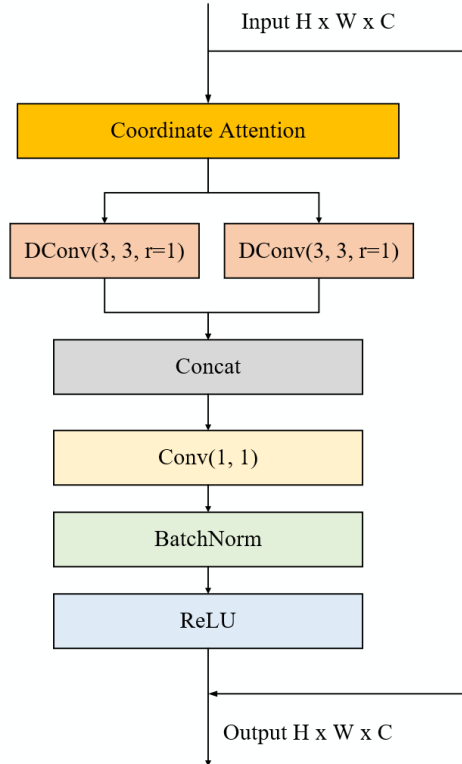


Figure 2: Architecture of CFGM

**Motivation.** Although the FPN in CLRNet effectively merges multi-scale features, it relies on standard  $3 \times 3$  convolutions that process spatial dimensions isotropically. However, lane markings exhibit strong anisotropic characteristics—elongated horizontally and thin vertically. Furthermore,

ROI-Gather aggregates features from all channels equally, including redundant ones activated by background clutter. To address these limitations, CFGM is designed to explicitly model positional dependencies and expand receptive fields along the lane direction, as illustrated in Fig. 2.

**Coordinate Attention (CA).** Coordinate Attention is first applied to each FPN output. Unlike channel attention mechanisms (e.g., SE) that squeeze spatial information globally, CA encodes horizontal and vertical coordinates via two one-dimensional global pooling operations:

$$\begin{aligned} z^h &= \frac{1}{W} \sum_{0 \leq i < W} x(h, i), \\ z^w &= \frac{1}{H} \sum_{0 \leq j < H} x(j, w). \end{aligned}$$

These coordinate-aware embeddings are concatenated, processed by a shared  $1 \times 1$  convolution, and split into two attention vectors  $g^h \in \mathbb{R}^{C \times H \times 1}$  and  $g^w \in \mathbb{R}^{C \times 1 \times W}$ . The recalibrated feature is obtained by:

$$x' = x \cdot \sigma(g^h) \cdot \sigma(g^w),$$

where  $\sigma(\cdot)$  denotes the sigmoid function. CA preserves precise positional information along the lane direction while suppressing irrelevant background responses.

**Dual Dilated Convolutions.** After CA recalibration, two parallel dilated convolution branches are employed to capture multi-scale context without resolution loss:

$$\begin{aligned} f_1 &= \text{Conv}_{3 \times 3}^{r=1}(x'), \\ f_3 &= \text{Conv}_{3 \times 3}^{r=3}(x'). \end{aligned}$$

The branch with  $r = 1$  preserves fine-grained local details, while the branch with  $r = 3$  expands the receptive field to  $7 \times 7$ , capturing long-range dependencies essential for curved lanes and distant vanishing points. The two branches are concatenated and fused:

$$f = \text{Conv}_{1 \times 1}([f_1; f_3]).$$

**Residual Connection.** To preserve the original feature flow and stabilize training, the fused feature is added back to the input via a residual connection:

$$y = \text{ReLU}(\text{BN}(f) + x).$$

Each CFGM introduces negligible parameters yet significantly enhances feature discriminability.

### 3.3. Geometry Consistency Loss

**Motivation.** The training loss of CLNet comprises a classification loss (Focal loss) and a regression loss (Line IoU loss + Smooth L1). While these losses ensure point-wise matching accuracy, they do not explicitly constrain the geometric smoothness of the predicted curve. Consequently, under occlusions or worn lane markings, the network tends to produce fragmented or oscillating predictions.

**Formulation.** A Geometry Consistency Loss (GCL) based on second-order difference constraints is proposed. For a predicted lane prior with  $N$  sampled points  $p_i = (x_i, y_i)$ , the first-order differences are:

$$d_i = p_{i+1} - p_i, \quad i = 1, \dots, N - 1.$$

The second-order difference, which measures the curvature variation between consecutive

segments, is defined as:

$$\Delta^2 p_i = d_{i+1} - d_i = p_{i+2} - 2p_{i+1} + p_i.$$

The Geometry Consistency Loss penalizes large second-order variations, enforcing smooth transitions:

$$\mathcal{L}_{\text{GCL}} = \frac{1}{N-2} \sum_{i=1}^{N-2} \|\Delta^2 p_i\|_1.$$

**Total Loss.** The overall training objective is formulated as:

$$\mathcal{L}_{\text{total}} = w_{\text{cls}} \mathcal{L}_{\text{cls}} + w_{\text{iou}} \mathcal{L}_{\text{LineIoU}} + w_{\text{xytl}} \mathcal{L}_{\text{xytl}} + \lambda \mathcal{L}_{\text{GCL}},$$

where  $\mathcal{L}_{\text{cls}}$  denotes the Focal loss,  $\mathcal{L}_{\text{LineIoU}}$  the Line IoU loss,  $\mathcal{L}_{\text{xytl}}$  the Smooth L1 loss for start point coordinates, angle, and length, and  $\lambda$  is empirically set to 0.1 in all experiments.

### 3.4. Training and Inference Details

**Positive Sample Selection.** Following CLNet, the SimOTA dynamic label assignment strategy is adopted. Each ground-truth lane is matched with the top-k predicted lanes based on an assignment cost:

$$C_{\text{assign}} = w_{\text{sim}} C_{\text{sim}} + w_{\text{cls}} C_{\text{cls}},$$

where  $C_{\text{sim}} = C_{\text{dis}} \cdot C_{\text{xy}} \cdot C_{\theta}$  integrates pixel distance, start point distance, and angle difference.

**Inference.** During inference, low-confidence lane priors are filtered using a classification score threshold of 0.6. Non-Maximum Suppression (NMS) with a line IoU threshold of 0.5 is applied to eliminate redundant predictions.

## 4. Experiments

### 4.1. Datasets and Evaluation Metrics

We evaluate our method on two widely used lane detection benchmarks: CULane [7] and TuSimple [15]. CULane [7] is a large-scale dataset designed for complex urban driving scenarios. It contains 100,000 images ( $1640 \times 590$  resolution) covering nine challenging categories, including crowded scenes, night, and crossroad conditions. The dataset provides rich annotations to benchmark lane detection performance under diverse real-world environments. TuSimple [15] focuses on highway driving scenes, providing 6,408 images ( $1280 \times 720$  resolution) split into 3,268 training, 358 validation, and 2,782 test images. Its clean annotations and controlled scenarios make it a standard benchmark for evaluating lane detection accuracy in highway settings.

The F1 score is adopted as the primary evaluation metric, which is computed based on the intersection-over-union (IoU) between predicted and ground-truth lanes. For CULane, the F1 score is reported at IoU thresholds of 0.5 and 0.75, along with category-specific F1 scores for nine scenarios. For TuSimple, the accuracy, false positive (FP), and false negative (FN) rates are additionally reported.

### 4.2. Implementation Details

DLA34 and ResNet are employed as the pre-trained backbone networks. The CULane dataset is trained for 15 epochs, while TuSimple is trained for 70 epochs. The AdamW optimizer is adopted

with an initial learning rate of  $1e-3$ , and a cosine annealing learning rate schedule is applied with the power parameter set to 0.9. The network is implemented based on PyTorch, and all experiments are conducted on a single GPU. The number of lane prior points is set to  $N = 72$ , and the number of sampled points is  $N_p = 36$ .

#### 4.3. Comparison with State-of-the-Art Methods

Performance on CULane. As shown in Table 1, the proposed method achieves competitive performance on the CULane benchmark. The overall F1@50 reaches 80.88%, surpassing CLRNNet by 0.41%. Among various challenging scenarios, the most significant improvement is observed on the Curve category (74.82%, +0.69%), which validates the effectiveness of the dual-branch dilated convolutions in modeling long-range contextual dependencies for curved lanes. The Crowded category achieves 80.20% (+0.61%), demonstrating that the Geometry Consistency Loss plays a critical role in maintaining lane continuity under occlusion conditions. The Night category reaches 75.91% (+0.54%), indicating that Coordinate Attention effectively enhances positional awareness in low-light environments. Additionally, the False Positive count on the Cross category is reduced from 1155 to 1115 (a decrease of 40), reflecting the improved false detection suppression capability at complex intersections.

Performance on TuSimple. As shown in Table 2, on the nearly saturated TuSimple dataset, the proposed method still achieves  $F1 = 97.95\%$  (+0.06%) and  $Acc = 96.89\%$  (+0.05%). More importantly, the False Positive rate is reduced from 2.28% to 2.14% (a decrease of 0.14%), and the False Negative rate is reduced from 1.92% to 1.83% (a decrease of 0.09%). These results indicate that the proposed method effectively reduces both false detections and missed detections while maintaining high accuracy, validating its robustness in clean highway scenarios.

Table 1: State-of-the-art results on CULane

| Method          | Backbone     | mF1          | F1@50        | F1@75        | Normal       | Crowded      | Dazzle       | Shadow       | No Line      | Arrow        | Curve        | Cross       | Night        |
|-----------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|-------------|--------------|
| RESA            | ResNet50     | 47.86        | 75.30        | 53.39        | 92.10        | 73.10        | 69.20        | 72.80        | 47.70        | 88.30        | 70.30        | 1503        | 69.90        |
| LaneATT         | ResNet122    | 51.48        | 77.02        | 57.50        | 91.74        | 76.16        | 69.47        | 76.31        | 50.46        | 86.29        | 64.05        | 1264        | 70.81        |
| CondLane        | ResNet101    | 54.83        | 79.48        | 61.23        | 93.47        | 77.44        | 70.93        | 80.91        | 54.13        | 90.16        | 75.21        | 1201        | 74.80        |
| CLRNNet         | DLA34        | 55.64        | 80.47        | 62.78        | 93.73        | 79.59        | 75.30        | 82.51        | 54.58        | 90.62        | 74.13        | 1155        | 75.37        |
| <b>Proposed</b> | <b>DLA34</b> | <b>56.08</b> | <b>80.88</b> | <b>63.51</b> | <b>93.99</b> | <b>80.20</b> | <b>75.74</b> | <b>82.91</b> | <b>54.98</b> | <b>90.92</b> | <b>74.82</b> | <b>1115</b> | <b>75.91</b> |

Note: The Cross column reports False Positive counts (lower is better); all other columns report F1 scores (%).

Table 2: State-of-the-art results on TuSimple

| Method      | Backbone  | F1(%)        | Acc(%)       | FP(%)       | FN(%)       |
|-------------|-----------|--------------|--------------|-------------|-------------|
| RESA        | ResNet34  | 96.93        | 96.82        | 3.63        | 2.48        |
| UFLD        | ResNet34  | 88.02        | 95.86        | 18.91       | 3.75        |
| LaneATT     | ResNet122 | 96.06        | 96.10        | 5.64        | 2.17        |
| FOLOLane    | ERFNet    | 96.59        | <b>96.92</b> | 4.47        | 2.28        |
| CondLaneNet | ResNet34  | 96.98        | 95.37        | 2.20        | 3.82        |
| CLRNNet     | ResNet18  | 97.89        | 96.84        | 2.28        | 1.92        |
| Proposed    | ResNet18  | <b>97.95</b> | 96.89        | <b>2.14</b> | <b>1.83</b> |

#### 4.4. Ablation Study

To verify the effectiveness of each proposed component, ablation experiments are conducted on the CULane dataset, and the results are presented in Table 3.

Overall ablation analysis. Taking CLRNNet as the baseline ( $F1 = 79.58\%$ ), the introduction of

CFGM improves the overall F1 to 80.00% (+0.42%). The Curve, Night, and Crowded scenarios are improved by 0.62%, 0.45%, and 0.43%, respectively, indicating that CFGM effectively fuses multi-scale contextual information and enhances feature representation in complex scenarios. Upon further introducing GCL, the overall F1 reaches 80.25% (a cumulative improvement of +0.67%). The Crowded and Night scenarios are further improved by 0.18% and 0.09%, respectively, while the Curve scenario gains only 0.07%, suggesting that the geometric consistency constraint primarily improves lane continuity under occlusion and partial missing conditions.

CFGM analysis. CFGM achieves the largest gain on the Curve scenario (+0.62%), validating the effectiveness of dual-branch dilated convolutions in modeling long-range contextual dependencies. The Night scenario is improved by 0.45%, demonstrating that Coordinate Attention enhances spatial feature expression under low-light conditions. The Crowded scenario is improved by 0.43%, further indicating that multi-scale feature fusion helps alleviate feature degradation caused by occlusion.

GCL analysis. GCL achieves the largest gain on the Crowded scenario (+0.18%), indicating that the geometric consistency constraint effectively suppresses prediction fractures caused by occlusion and generates smoother and more continuous lane curves. The Night scenario is further improved by 0.09%, suggesting that geometric priors help maintain correct lane topology under complex environments. The gain on the Curve scenario is relatively small (0.07%), indicating that performance improvements in this scenario are mainly attributed to CFGM, while GCL primarily serves as a geometric optimizer for the prediction results.

Table 3: Effects of each component in our method. Results are reported on CULane

| Baseline | CFGM | GCL | Crowded | Curve | Night | F1@50 |
|----------|------|-----|---------|-------|-------|-------|
| ✓        |      |     | 79.59   | 74.13 | 75.37 | 79.58 |
| ✓        | ✓    |     | 80.02   | 74.75 | 75.82 | 80.00 |
| ✓        | ✓    | ✓   | 80.20   | 74.82 | 75.91 | 80.25 |

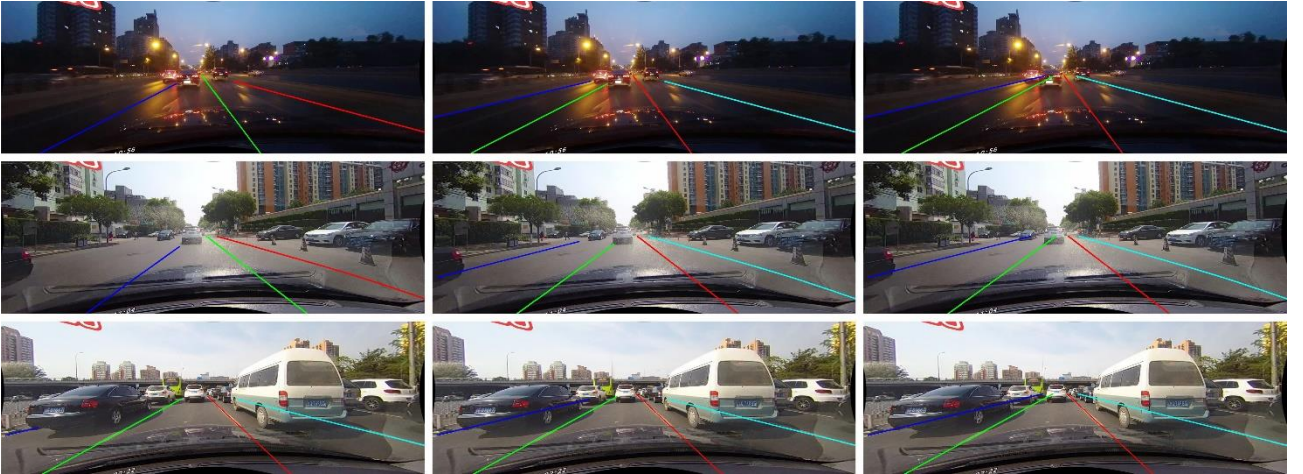


Figure 3: Visulization results on Culane dataset. Three representative challenging scenarios are selected for visualization, including the night, no-line, and crowded scenes (from top to bottom), where the three columns from left to right show the detection results of CLRNNet, the proposed method, and the ground truth, respectively.

## 4.5. Qualitative Results

### 4.5.1. Results on CULane Dataset

As illustrated in Fig.3, we present the qualitative comparison between the baseline CLRNet and the proposed method on the CULane dataset. Three representative challenging scenarios are selected for visualization, including the night, no-line, and crowded scenes (from top to bottom), where the three columns from left to right show the detection results of CLRNet, the proposed method, and the ground truth, respectively. It can be clearly observed that the baseline suffers from severe missing detection in all three scenarios, while the proposed method consistently produces complete and accurate lane predictions that are highly consistent with the ground truth. Specifically, in the night scene (the first row), CLRNet fails to detect the lanes due to the poor illumination condition, whereas our method still yields reliable predictions, which benefits from the context feature guidance of the CFGM that enhances lane representations under adverse lighting. In the no-line and crowded scenes (the second and third rows), where the lane markings are worn out or heavily occluded by surrounding vehicles, CLRNet generates fragmented predictions with obvious missing segments, while the proposed method can still infer continuous lane structures, demonstrating the effectiveness of the geometry continuity constraint imposed by the GCL. These qualitative results further verify the robustness and superiority of the proposed method in challenging scenarios.

### 4.5.2. Results on Tusimple Dataset

Fig. 4 further shows the qualitative comparison on the TuSimple dataset, which is collected in highway scenarios with stable illumination and clear lane markings. Following the same layout as Fig. 3, the three columns from left to right present the detection results of CLRNet, the proposed method, and the ground truth, respectively. Although the highway scenes in TuSimple are relatively simple, accurately detecting the distant lanes near the vanishing area remains challenging, since these lanes occupy only a few pixels and are easily confused with the background. As shown in the first column, CLRNet produces incomplete predictions for the distant lanes, where obvious breakpoints and missing segments can be observed. In contrast, benefiting from the context feature guidance of the CFGM, the proposed method effectively enhances the weak features of distant lanes and generates more complete and continuous predictions. Moreover, for the curved and occluded lanes caused by the preceding vehicles, our method maintains smooth lane structures that are highly consistent with the ground truth, which further demonstrates the effectiveness of the geometry continuity constraint. These results indicate that the proposed method generalizes well to different scenarios and achieves robust lane detection performance on both datasets.

## 5. Conclusions

In this paper, we proposed **CFG-CLRNet** for robust vision-based lane detection in complex road scenes. Built upon CLRNet, the proposed framework enhances lane detection from both contextual feature modeling and geometric structure optimization. Specifically, the Context Feature Guidance Module integrates coordinate attention with dual-branch atrous convolution to strengthen multi-scale contextual interaction while preserving position-sensitive lane information. Meanwhile, the Geometric Consistency Loss imposes a second-order smoothness constraint on lane points, encouraging more continuous and geometrically consistent predictions under occlusion, illumination variation, and incomplete lane markings. Extensive experiments on CULane and TuSimple demonstrate that CFG-CLRNet consistently improves baseline performance and shows

stronger robustness in challenging scenarios.

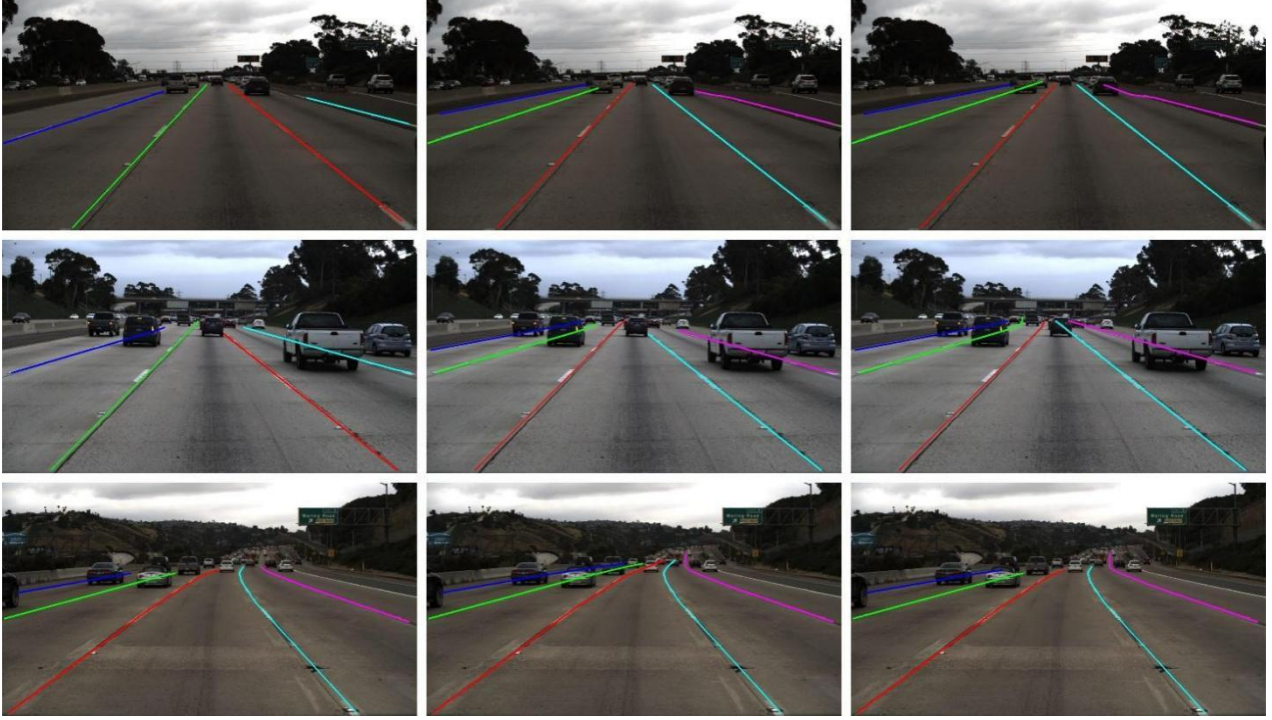


Figure 4: Visualization results on Tusimple dataset. Three representative challenging scenarios are selected for visualization, including the night, no-line, and crowded scenes (from top to bottom), where the three columns from left to right show the detection results of CLRNNet, the proposed method, and the ground truth, respectively

## Acknowledgements

This paper is supported by the school-level industry-education integration course Introduction to “Civil Aviation”.

## References

- [1] Xu, H., Wang, S., Cai, X., Zhang, W., Liang, X., & Li, Z. (2020). Curvelane-nas: Unifying lane-sensitive architecture search and adaptive point blending. In *European Conference on Computer Vision*, 689-704.
- [2] Zheng, T., Fang, H., Zhang, Y., Tang, W., Yang, Z., Liu, H., & Cai, D. (2021). Resa: Recurrent feature-shift aggregator for lane detection. In *Proceedings of the AAAI conference on artificial intelligence*, 3547-3554.
- [3] Yoo, S., Lee, H. S., Myeong, H., Yun, S., Park, H., Cho, J., & Kim, D. H. (2020). End-to-end lane marker detection via row-wise classification. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition workshops*, 1006-1007.
- [4] Tabelini, L., Berriel, R., Paixao, T. M., Badue, C., De Souza, A. F., & Oliveira-Santos, T. (2021). Polylandenet: Lane estimation via deep polynomial regression. In *2020 25th international conference on pattern recognition (ICPR)*, 6150-6156.
- [5] Li, X., Li, J., Hu, X., & Yang, J. (2019). Line-cnn: End-to-end traffic line detection with line proposal unit. *IEEE Transactions on Intelligent Transportation Systems*, 21(1), 248-258.
- [6] Zheng, T., Huang, Y., Liu, Y., Tang, W., Yang, Z., Cai, D., & He, X. (2022). Clrnet: Cross layer refinement network for lane detection. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 898-907.
- [7] Pan, X., Shi, J., Luo, P., Wang, X., & Tang, X. (2018). Spatial as deep: Spatial cnn for traffic scene understanding. In *Proceedings of the AAAI conference on artificial intelligence*.
- [8] Liu, R., Yuan, Z., Liu, T., & Xiong, Z. (2021). End-to-end lane shape prediction with transformers. In *Proceedings of the IEEE/CVF winter conference on applications of computer vision*, 3694-3702.
- [9] Feng, Z., Guo, S., Tan, X., Xu, K., Wang, M., & Ma, L. (2022). Rethinking efficient lane detection via curve

- modeling. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 17062-17070.
- [10] Liu, L., Chen, X., Zhu, S., & Tan, P. (2021). Condlanenet: a top-to-down lane detection framework based on conditional convolution. In *Proceedings of the IEEE/CVF international conference on computer vision*, 3773-3782.
- [11] Tabelini, L., Berriel, R., Paixao, T. M., Badue, C., De Souza, A. F., & Oliveira-Santos, T. (2021). Keep your eyes on the lane: Real-time attention-guided lane detection. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 294-302.
- [12] Qin, Z., Wang, H., & Li, X. (2020). Ultra fast structure-aware deep lane detection. In *European conference on computer vision*, 276-291.
- [13] Honda, H., & Uchida, Y. (2024). Clrernet: improving confidence of lane detection with laneiou. In *Proceedings of the IEEE/CVF winter conference on applications of computer vision*, 1176-1185.
- [14] Su, J., Chen, C., Zhang, K., Luo, J., Wei, X., & Wei, X. (2021). Structure guided lane detection. *arXiv preprint arXiv:2105.05403*.
- [15] TuSimple: Tusimple benchmark. <https://github.com/TuSimple/tusimple-benchmark/> (2020).