



ARTICLE

Direction-Curvature Aware Feature Integration for Robust Lane Detection

Ahtisham Waheed¹, Yunfie Yin^{1,*}, Abu Fatema Mohammad Abdun Noor², Md Imam Ahasan¹,
Kah Ong Michael Goh^{3,*}, S. M. Hasan Mahmud^{2,*} and Umar Rashid⁴

¹College of Computer Science, Chongqing University, Chongqing, China

²Department of Software Engineering, Daffodil International University, Dhaka, Bangladesh

³Center for Image and Vision Computing, COE for Artificial Intelligence, Faculty of Information Science & Technology, Multimedia University, Jalan Ayer Keroh Lama, Melaka, Malaysia

⁴Department of Computer Science, COMSATS University Islamabad, Islamabad, Pakistan

*Corresponding Authors: Yunfie Yin Email: yinyunfei@cqu.edu.cn; Kah Ong Michael Goh. Email: michael.goh@mmu.edu.my; S. M. Hasan Mahmud. Email: drhasan.swe@diu.edu.bd

Received: 04 April 2026; Accepted: 09 May 2026; Published: 13 August 2026

ABSTRACT: Robust lane detection is a fundamental perception task for autonomous driving and Advanced Driver Assistance Systems. However, it remains challenging in real-world environments due to degraded lane markings, complex road topologies, occlusions, and adverse illumination conditions. This work aims to improve lane detection robustness by explicitly modeling lane geometric properties while preserving end-to-end efficiency. We propose a direction-curvature aware lane detection framework that integrates a novel Direction-Curvature Aware (DCA) attention module into an anchor-based architecture. The DCA module enables tangent-aligned feature aggregation guided by learned direction fields and curvature-consistent attention. In addition, we introduce a direction-aware optimization objective termed Directional Lane IoU (DLIoU) to enforce directional consistency between predicted and ground-truth lanes during training. Extensive experiments on the CULane and TuSimple benchmarks demonstrate that the proposed method consistently outperforms state-of-the-art approaches, achieving notable improvements on CULane under challenging conditions including occlusion, strong curvature, shadows, and night-time scenes, and consistent gains on TuSimple, while maintaining competitive inference speed. These results confirm that explicitly incorporating direction and curvature information into both feature learning and optimization leads to more accurate and robust lane detection in complex driving environments.

KEYWORDS: Lane detection; deep learning; attention mechanism; curvature-aware modeling; anchor-based detection; autonomous driving; computer vision

1 Introduction

Lane detection estimates the geometric structure and spatial layout of lane markings and is a core perception module for autonomous driving and Advanced Driver Assistance Systems (ADAS). Reliable lane perception is essential for lane keeping and safety-critical planning and control in dynamic traffic environments [1]. Despite steady progress, robust lane detection remains difficult in real-world settings: markings are often degraded, partially missing, or visually ambiguous due to low illumination (tunnels), dense-traffic occlusions, complex intersection topologies, and appearance degradation from shadows, glare, and worn road surfaces [2–4]. Deep learning (DL), particularly Convolutional Neural Networks (CNNs), has significantly advanced lane detection accuracy and real-time performance [4–6]. However, robustness under

complex environments remains limited, largely due to insufficient joint modeling of lane geometry and long-range contextual dependencies. Specifically, segmentation-based methods are sensitive to occlusions and illumination changes due to their reliance on dense pixel-wise predictions, while anchor-based approaches struggle with complex topologies and strong curvature owing to their dependence on fixed geometric priors, collectively highlighting the need for explicit geometric modeling under adverse visual conditions. Early DL approaches treated lane detection as semantic segmentation [7,8] and then grouped pixels into lane instances via post-processing such as clustering or curve fitting [9,10]. While effective, these pipelines are not fully end-to-end and incur notable computational overhead. Anchor-based formulations [11–13] recast lanes as detection targets, yet their dependence on hand-crafted narrow anchors [14,15] and non-maximum suppression [16] still impedes truly end-to-end learning. As illustrated in Fig. 1, existing lane detection methods exhibit a clear trade-off between inference speed and detection performance on the CULane benchmark, whereas DCANet(our) achieves a favorable balance by lying on the Pareto frontier with superior Normal F1-score and competitive real-time efficiency.

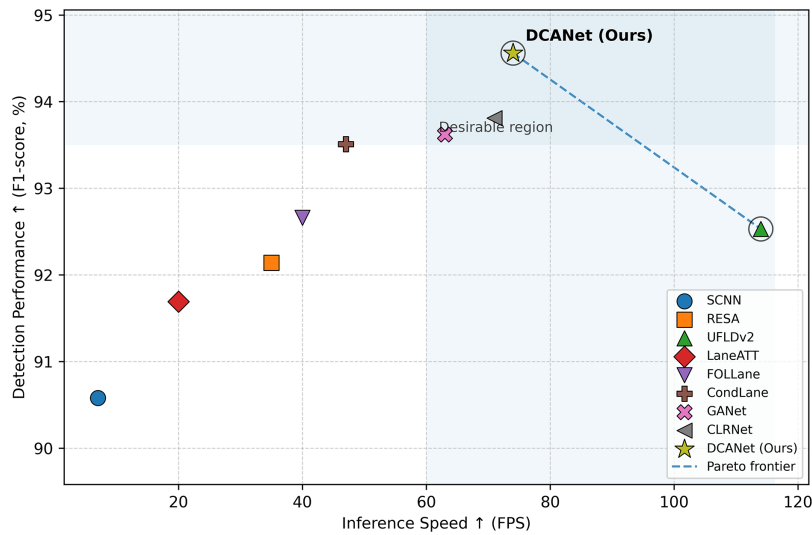


Figure 1: Speed-accuracy trade-off on the CULane benchmark under the Normal scenario. Each point represents a lane detection method evaluated by inference speed (FPS) and Normal F1-score. The Pareto frontier denotes the set of non-dominated methods that achieve the optimal trade-off between speed and accuracy, where the proposed DCANet lies on the frontier and attains superior detection performance with competitive real-time efficiency.

To simplify representation, PolyLaneNet [17] regresses cubic polynomial coefficients for lane lines, but this often decouples lane geometry from global road context and degrades performance in dense traffic or damaged road scenes. Context modeling has been strengthened via spatial message passing [4,18,19] and by exploiting additional scene annotations [20,21], at the cost of higher computation and annotation burden. LSTR [22] adopts a Transformer based architecture [23] with camera-specific curve formulations to capture long-range dependencies, but complex curve modeling and transformer training challenges can limit robustness in difficult scenes [4]. Overall, existing approaches motivate lane detection frameworks that explicitly encode geometric continuity and curvature while efficiently integrating context in an end-to-end manner.

We propose a direction-curvature aware lane detection framework that explicitly models lane geometry and integrates contextual information within an end-to-end architecture. Our approach leverages curvature-sensitive representations to preserve geometric continuity under complex topologies and adverse visual conditions. By jointly enforcing directional consistency and curvature-aware spatial relationships, the

method improves robustness to occlusions, illumination variation, and degraded markings without excessive computational overhead. Our contributions are:

- A direction-curvature aware feature integration framework that explicitly encodes lane geometric continuity for robustness under complex topologies and challenging visual conditions.
- A geometry-aware feature fusion strategy that combines local directional cues with global contextual information in an end-to-end trainable architecture.
- Extensive experiments on public benchmarks showing consistent improvements over state-of-the-art methods, especially under occlusions, illumination changes, and degraded lane markings.

2 Related Work

2.1 Segmentation-Based Methods

Segmentation-based methods [4,19,24–28] treat lane detection as pixel-wise prediction and are commonly grouped into semantic and instance segmentation. Semantic approaches detect lane pixels and then assemble them into lane instances: LaneAF [24] predicts segmentation and affinity fields for pixel association, FOLLane [25] refines lane points and links them to reconstruct lanes, and LaneNet [26] uses embedding learning with clustering to separate instances. Instance segmentation explicitly models each lane as an object: SCNN [4] and RESA [19] propagate context via spatial message passing; CurveLane-NAS [27] applies Neural Architecture Search for geometry-aware designs; and CondLaneNet [28] and CondLSTR [29] use dynamic kernel generation to extract lane instances from feature maps. While accurate, segmentation-based pipelines can be sensitive to occlusions and illumination changes, reflecting difficulty in jointly preserving fine-grained geometric continuity and global semantics.

2.2 Row-Wise Classification Methods

Row-wise classification improves efficiency by avoiding dense pixel-wise maps [30]. The image is discretized into horizontal grids, and lane positions are predicted per row under predefined priors, consistency constraints then reconstruct lane curves [10]. However, these constraints often induce non-trivial post-processing and reduce adaptability to diverse lane topologies. CondLaneNet [28] improves instance-level precision by additionally predicting vertical extents and refining offsets, but this increases training and post-processing complexity.

2.3 Anchor-Based Methods

Anchor-based approaches [5,6,12–15,30–34] represent lanes with predefined anchors, typically as line-anchor-based or row-anchor-based models. Line-CNN [12] introduced line anchors with regression refinement, LaneATT [13] aggregates global context via anchor-based attention; and SGNet [31] uses vanishing-point-guided anchor generation and structural constraints. Row-anchor-based methods such as UFLD [30] achieve high efficiency with lightweight backbones but can sacrifice accuracy, UFLDv2 [15] jointly models row and column information to better capture curvature; and LaneFormer [5] employs Transformer attention to model long-range spatial relations. Fixed anchors still struggle with complex topology and varying curvature: CLRNet [14] introduces learnable anchors and ROIgather for multi-level context fusion, ADNet [6] improves anchor initialization dynamically, and CLRerNet [32] proposes LaneIoU by incorporating local lane orientation for better confidence estimation and accuracy.

3 Methods

3.1 Overall Architecture

The proposed framework adopts an anchor-based lane detection paradigm and is designed as an end-to-end trainable architecture. It consists of four main components: a DLA-34 backbone for hierarchical feature extraction, a Direction-Curvature Aware (DCA) attention module for directional feature refinement, a context-aware feature pyramid neck for multi-scale feature aggregation, and an anchor-based lane detection head for final lane prediction. As illustrated in Fig. 2, the DCA module is introduced as the only backbone-level modification, preserving the baseline architecture while enhancing geometric representation. Given an input image $\mathbf{x} \in \mathbb{R}^{3 \times H \times W}$, the backbone extracts multi-scale feature maps as

$$\{\mathbf{C}_2, \mathbf{C}_3, \mathbf{C}_4, \mathbf{C}_5\} = \mathcal{B}(\mathbf{x}), \quad (1)$$

where $\mathcal{B}(\cdot)$ denotes the backbone network and $\mathbf{C}_i \in \mathbb{R}^{C_i \times \frac{H}{2^i} \times \frac{W}{2^i}}$ represents the feature map at stride 2^i . Lower-level features capture fine-grained spatial details, while higher-level features encode richer semantic and contextual information. To explicitly model directional continuity and curvature, the DCA module is applied exclusively to the deepest feature map $\mathbf{C}_5$, yielding

$$\mathbf{C}'_5 = \mathcal{D}(\mathbf{C}_5), \quad (2)$$

where $\mathcal{D}(\cdot)$ denotes the DCA operation. This design preserves both the spatial resolution and channel dimensionality of $\mathbf{C}_5$, enabling seamless integration with the baseline pipeline. The refined feature $\mathbf{C}'_5$ and the remaining backbone outputs are then fused by a context-aware feature pyramid neck:

$$\{\mathbf{P}_i\}_{i=2}^5 = \mathcal{N}(\mathbf{C}_2, \mathbf{C}_3, \mathbf{C}_4, \mathbf{C}'_5), \quad (3)$$

where $\mathcal{N}(\cdot)$ denotes the neck network and $\mathbf{P}_i$ are the resulting pyramid features. Finally, the anchor-based detection head operates on the pyramid features to predict lane candidates:

$$\{\mathbf{s}, \mathbf{r}, \mathbf{e}\} = \mathcal{H}(\mathbf{P}_2, \mathbf{P}_3, \mathbf{P}_4, \mathbf{P}_5), \quad (4)$$

where $\mathcal{H}(\cdot)$ denotes the detection head, and $\mathbf{s}$, $\mathbf{r}$, and $\mathbf{e}$ correspond to classification, regression, and existence predictions, respectively. By introducing the DCA module as the sole architectural modification, the proposed framework ensures fair comparison with existing anchor-based baselines while effectively incorporating Direction-curvature aware feature representations. The Direction-Curvature Aware (DCA) attention module is applied exclusively to the deepest backbone feature map $\mathbf{C}_5$ to ensure reliable geometric modeling with minimal computational overhead. In the DLA-34 backbone, shallow features $\mathbf{C}_2 \in \mathbb{R}^{64 \times \frac{H}{4} \times \frac{W}{4}}$ and $\mathbf{C}_3 \in \mathbb{R}^{128 \times \frac{H}{8} \times \frac{W}{8}}$ predominantly capture low-level appearance cues, while $\mathbf{C}_4$ encodes intermediate structures. In contrast, $\mathbf{C}_5 \in \mathbb{R}^{512 \times \frac{H}{32} \times \frac{W}{32}}$ provides semantically stable representations over a large receptive field, which is critical for modeling lane geometry. The per-pixel direction field $\mathbf{V}(p)$ predicted by DCA is designed to approximate the local tangent of a lane,

$$\mathbf{V}(p) \approx \frac{\partial \mathcal{L}(s)}{\partial s},$$

where $\mathcal{L}(s)$ denotes a parametric lane curve. Such geometric quantities require global context and are difficult to infer from shallow, noise-sensitive features. Applying DCA at $\mathbf{C}_5$ enables tangent-aligned sampling,

$$p_k = p + \delta_k s \mathbf{V}(p),$$

to operate on coherent feature representations, thereby promoting consistent aggregation along lane trajectories. Since the computational complexity of DCA scales with spatial resolution as $\mathcal{O}(HWKC)$, restricting it to the lowest-resolution feature map preserves efficiency while allowing direction-aware information to propagate through subsequent multi-scale fusion.

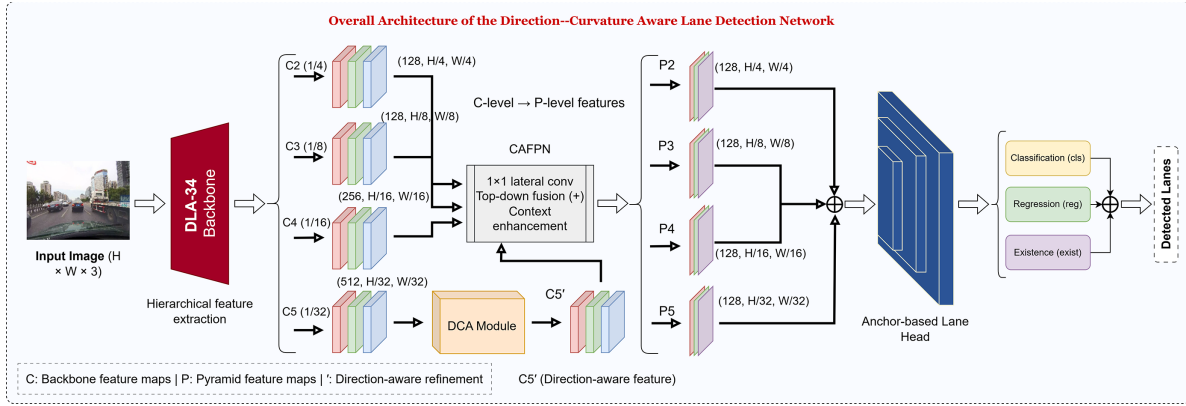


Figure 2: Overall architecture of the proposed method. A DLA-34 backbone extracts multi-scale features, where the proposed DCA module refines the deepest feature map. The refined features are fused by a context-aware feature pyramid and processed by an anchor-based lane detection head to produce final lane predictions.

3.2 Direction-Curvature Aware Attention

Standard convolutional and self-attention operators aggregate features within spatially symmetric neighborhoods, which is ill-suited for modeling such anisotropic geometry. To address this limitation, we propose the *Direction-Curvature Aware* (DCA) attention module, which performs tangent-aligned feature aggregation guided by a learned direction field. Let $C_5 \in \mathbb{R}^{C \times H \times W}$ denote the deepest backbone feature map. DCA first estimates a dense per-pixel direction field that approximates the local tangent orientation of lane markings. The direction field is predicted using a lightweight estimation network composed of two consecutive 1×1 convolutional layers with a ReLU activation in between:

$$\mathbf{v}_{\text{raw}} = \text{Conv}_{1 \times 1}(\sigma(\text{Conv}_{1 \times 1}(C_5))), \quad (5)$$

where $\sigma(\cdot)$ denotes the ReLU function. The output $\mathbf{v}_{\text{raw}} \in \mathbb{R}^{2 \times H \times W}$ represents unnormalized two-dimensional direction vectors. These vectors are normalized at each spatial location to ensure scale invariance and numerical stability:

$$\mathbf{V} = \frac{\mathbf{v}_{\text{raw}}}{\|\mathbf{v}_{\text{raw}}\|_2 + \epsilon}, \quad (6)$$

where $\mathbf{v}_{\text{raw}}$ represents the unnormalized local direction vectors predicted from the feature map, $\mathbf{V}$ denotes the corresponding unit direction field, and ϵ is a small constant for numerical stability. The normalization ensures that the direction vectors encode only orientation information, independent of magnitude, thereby providing a consistent tangent direction at each spatial location. Based on this unit direction field, the DCA module performs tangent-aligned feature sampling to aggregate contextual information along the predicted lane direction. Specifically, for each spatial location $\mathbf{p} = (x, y)$, a set of K sampling points is defined as

$$\mathbf{p}_k = \mathbf{p} + \Delta_k \mathbf{V}(\mathbf{p}), \quad k = 1, \dots, K, \quad (7)$$

where $\{\Delta_k\}$ are evenly spaced scalar offsets. Feature values at these locations are obtained using bilinear interpolation, producing a set of sampled tokens $\{\mathbf{t}_k \in \mathbb{R}^C\}$. This sampling strategy aggregates features along the lane tangent while suppressing interference from orthogonal directions. Although the sampling offsets in Eq. (7) follow a locally linear trajectory defined by the tangent direction $\mathbf{V}(\mathbf{p})$, curvature awareness in DCA arises from two complementary mechanisms. First, $\mathbf{V}(\mathbf{p})$ is predicted densely and independently at each spatial location, so the direction field naturally adapts to curved lane geometries across the feature map. The local linear approximation remains geometrically valid within the small neighborhood spanned by the sampling offsets, particularly at the stride-32 resolution of $\mathbf{C}_5$, where each spatial unit corresponds to a 32-pixel region in the input image, effectively limiting the physical extent of the sampling window. Second, the curvature gate g_k in Eq. (9) explicitly enforces curvature consistency by measuring the cosine similarity between the direction vectors at the center location $\mathbf{p}$ and each sampled point $\mathbf{p}_k$. Under strong curvature, sampled points whose local direction deviates significantly from $\mathbf{V}(\mathbf{p})$ receive suppressed attention weights, preventing geometrically inconsistent features from corrupting the aggregated representation. Together, these two mechanisms enable DCA to remain geometrically coherent under both straight and curved lane conditions without requiring explicit curved sampling paths.

To adaptively fuse the sampled features, DCA applies a local attention mechanism centered at the reference location. The center token and sampled tokens are projected into query, key, and value embeddings using linear transformations:

$$\mathbf{q} = \mathbf{W}_q \mathbf{t}_0, \quad \mathbf{k}_k = \mathbf{W}_k \mathbf{t}_k, \quad \mathbf{v}_k = \mathbf{W}_v \mathbf{t}_k \quad (8)$$

where $\mathbf{W}_q$, $\mathbf{W}_k$, and $\mathbf{W}_v$ are learnable weight matrices. To explicitly enforce curvature consistency, a directional gating term is introduced based on the cosine similarity between the direction vectors at the center and sampled locations:

$$g_k = \sigma(\alpha \mathbf{V}(\mathbf{p})^\top \mathbf{V}(\mathbf{p}_k) + \beta), \quad (9)$$

where α and β are learnable scalar parameters. This gate suppresses contributions from samples whose local direction deviates significantly from the center direction. The attention weights are then computed as

$$a_k = \text{softmax}(\mathbf{q}^\top \mathbf{k}_k + \log g_k), \quad (10)$$

which biases the aggregation toward directionally consistent features. The output of the attention operation is given by

$$\mathbf{y} = \sum_{k=1}^K a_k \mathbf{v}_k. \quad (11)$$

The aggregated feature $\mathbf{y}$ is projected back to the original channel dimension using a 1×1 convolution, followed by a depthwise 3×3 convolution for local smoothing. A residual connection is applied to preserve the original representation:

$$\mathbf{C}'_5 = \mathbf{C}_5 + \mathcal{P}(\mathbf{y}), \quad (12)$$

where $\mathcal{P}(\cdot)$ denotes the projection and smoothing operation. By combining direction field estimation, tangent-aligned sampling, and curvature-gated attention (Fig. 3), the DCA module enables explicit direction-curvature aware feature integration. This design improves robustness to curved and occluded lanes while remaining lightweight and fully compatible with standard convolutional backbones.

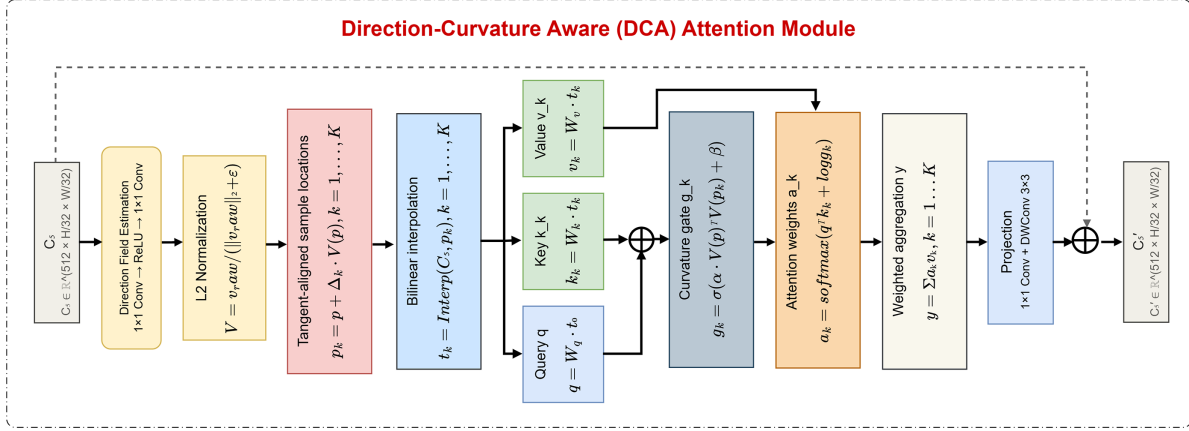


Figure 3: Internal architecture of the proposed Direction-Curvature Aware (DCA) module. The module estimates a unit direction field V from C_5 , samples $K = 9$ tangent-aligned tokens, applies curvature-gated local attention to produce aggregated feature y , and adds it back to C_5 via a residual connection to yield the refined output C'_5 .

Direction Field Learning: The direction field $V(p)$ is learned implicitly through the end-to-end training objective without explicit ground-truth tangent supervision. The detection loss backpropagates through the tangent-aligned sampling and curvature-gated attention, encouraging the network to predict direction fields that are geometrically consistent with lane structures. Regarding sign ambiguity, since lane annotations in both CULane and TuSimple are represented as ordered sequences of points from top to bottom following the row-anchor convention, the network naturally converges to a consistent directional orientation aligned with this ordering. This implicit canonical direction effectively resolves the sign ambiguity without requiring additional supervision or explicit sign normalization.

3.3 Direction-Aware Lane Optimization

Anchor-based lane detectors are commonly trained using classification and regression losses that emphasize spatial overlap and point-wise accuracy. While effective for localization, these objectives do not explicitly enforce directional continuity along the lane trajectory. As a result, predictions with locally inconsistent orientation or distorted curvature may still achieve low regression error, especially under occlusion or strong perspective distortion. To address this limitation, we introduce a direction-aware optimization objective that enforces directional agreement between predicted and ground-truth lanes during training. Let a predicted lane be represented by an ordered sequence of points $\hat{L} = \{\hat{\mathbf{p}}_i\}_{i=1}^N$ and the corresponding ground-truth lane by $L = \{\mathbf{p}_i\}_{i=1}^N$, where points are sampled at fixed row anchors. Local tangent directions are estimated using finite differences:

$$\begin{aligned} \hat{\mathbf{d}}_i &= \frac{\hat{\mathbf{p}}_{i+1} - \hat{\mathbf{p}}_i}{\|\hat{\mathbf{p}}_{i+1} - \hat{\mathbf{p}}_i\|_2 + \epsilon}, \\ \mathbf{d}_i &= \frac{\mathbf{p}_{i+1} - \mathbf{p}_i}{\|\mathbf{p}_{i+1} - \mathbf{p}_i\|_2 + \epsilon}, \end{aligned} \quad (13)$$

where ϵ is a small constant for numerical stability. These normalized vectors encode the local tangent direction between adjacent anchor intervals. Based on the estimated directions, we define a *Directional Lane IoU (DLIoU)* objective that captures both spatial alignment and directional consistency. Directional similarity at each anchor interval is computed using cosine similarity:

$$s_i = \hat{\mathbf{d}}_i^\top \mathbf{d}_i. \quad (14)$$

The DLIoU loss is formulated as

$$\mathcal{L}_{\text{DLIoU}} = 1 - \frac{1}{N-1} \sum_{i=1}^{N-1} s_i \omega_i, \quad (15)$$

where $\omega_i \in \{0, 1\}$ is a binary visibility flag defined as $\omega_i = \mathbb{1}[\mathbf{p}_i \text{ and } \mathbf{p}_{i+1} \text{ are both annotated and visible}]$, which excludes anchor intervals where either endpoint is missing or occluded from the directional supervision. This formulation encourages accurate localization while enforcing smooth directional evolution along visible lane segments only. The direction-aware loss is integrated into the overall training objective together with standard classification, regression, and existence losses. The final loss is given by

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{cls}} + \mathcal{L}_{\text{reg}} + \mathcal{L}_{\text{exist}} + \lambda_d \mathcal{L}_{\text{DLIoU}}, \quad (16)$$

where $\mathcal{L}_{\text{cls}}$ denotes the anchor classification loss, $\mathcal{L}_{\text{reg}}$ denotes the geometric regression loss, and $\mathcal{L}_{\text{exist}}$ supervises lane presence. The scalar λ_d controls the contribution of the direction-aware term. By embedding directional supervision directly into the optimization process, DLIoU aligns naturally with the direction-sensitive representations learned by the proposed DCA module. The proposed DLIoU differs from existing geometric loss functions in several key aspects. Standard IoU losses [15] measure spatial overlap between predicted and ground-truth lane regions but do not explicitly enforce directional consistency along the lane trajectory. LaneIoU [33] incorporates local lane orientation primarily for overlap estimation and confidence scoring during inference, but does not directly supervise tangent direction alignment during training. In contrast, DLIoU explicitly enforces local directional consistency at each anchor interval by aligning tangent directions between predicted and ground-truth lane points via cosine similarity. This directional supervision is embedded directly into the training objective, resulting in more stable and geometrically coherent predictions for curved and fragmented lanes where spatial overlap alone is insufficient to capture lane geometry. This results in more stable predictions for curved and fragmented lanes. The overall training and inference procedure of DCANet is summarized in Algorithm 1. To further illustrate the effectiveness of the proposed direction-curvature aware modeling, Fig. 4 presents the relative F1-score improvements of DCANet over the strongest baseline (CLRNet) across different CULane scenarios.

Algorithm 1: DCANet training and inference formulation

Require: $x \in \mathbb{R}^{3 \times H \times W}$, anchor set $\mathcal{A}$, row anchors $\mathcal{R}$, loss weight λ_d

Ensure: Predicted lane set $\hat{\mathcal{L}}$

- 1: $\{C_2, C_3, C_4, C_5\} \leftarrow \mathcal{B}(x)$
 - 2: $C_5^{\text{dca}} \leftarrow \mathcal{D}(C_5)$
 - 3: $\{P_i\}_{i=2}^5 \leftarrow \mathcal{N}(C_2, C_3, C_4, C_5^{\text{dca}})$
 - 4: $(\mathbf{s}, \mathbf{r}, \mathbf{e}) \leftarrow \mathcal{H}(\{P_i\}_{i=2}^5, \mathcal{A})$
 - 5: **if** training **then**
 - 6: $\mathcal{M} \leftarrow \mathcal{T}(\mathbf{s}, \mathbf{r}, \mathbf{e}; \mathcal{A}, \mathcal{R}, \mathcal{Y}^*)$
 - 7: $\mathcal{J}_{\text{cls}} \leftarrow \mathcal{L}_{\text{cls}}(\mathbf{s}, \mathcal{M})$
 - 8: $\mathcal{J}_{\text{reg}} \leftarrow \mathcal{L}_{\text{reg}}(\mathbf{r}, \mathcal{M})$
 - 9: $\mathcal{J}_{\text{exist}} \leftarrow \mathcal{L}_{\text{exist}}(\mathbf{e}, \mathcal{M})$
 - 10: $\mathcal{J}_{\text{DLIoU}} \leftarrow \mathcal{L}_{\text{DLIoU}}(\mathbf{r}, \mathcal{M})$
-

(Continued)

Algorithm 1 (continued)

```

11:  $\mathcal{J} \leftarrow \mathcal{J}_{\text{cls}} + \mathcal{J}_{\text{reg}} + \mathcal{J}_{\text{exist}} + \lambda_d \mathcal{J}_{\text{DLIoU}}$ 
12:  $\Theta \leftarrow \Theta - \eta \nabla_{\Theta} \mathcal{J}$ 
13: else
14:    $\tilde{\mathcal{L}} \leftarrow \mathcal{D}_f(\mathbf{r}, \mathbf{e}; \mathcal{A}, \mathcal{R})$ 
15:    $\hat{\mathcal{L}} \leftarrow \mathcal{P}(\tilde{\mathcal{L}}, \gamma)$ 
16: end if
17: return  $\hat{\mathcal{L}}$ 

```

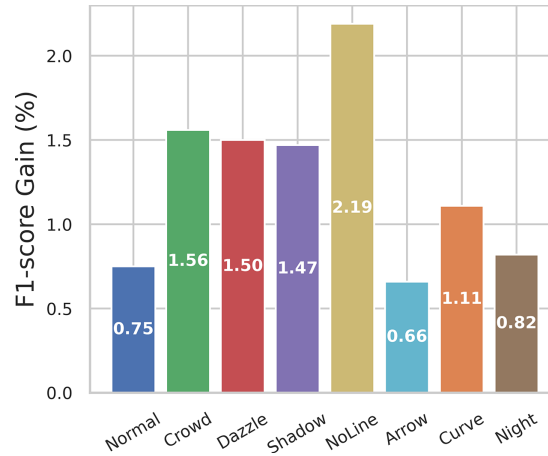


Figure 4: Relative F1-score gains of DCANet over the strongest baseline (CLRNet) across different CULane scenarios. The Cross scenario is excluded as it reports false positive counts instead of F1-score.

3.4 Implementation Details

Our implementation uses DLA-34 as the backbone network, with input resolutions of 320×800 for CULane and 368×640 for TuSimple. Standard data augmentations, including random horizontal flipping, photometric distortion, and geometric transformations, are applied consistently to both images and lane annotations during training. The Direction-Curvature Aware (DCA) attention module is inserted exclusively at the deepest backbone stage and operates on feature maps with 512 channels; it employs $K = 9$ tangent-aligned sampling points with a step size of 3.0 pixels and 4 attention heads with a per-head dimension of 64. Owing to its local and anisotropic design, the computational complexity of DCA scales linearly with spatial resolution as $\mathcal{O}(HWKC)$. All baseline hyperparameters, including anchor definitions, matching strategy, and learning rate schedule, are kept unchanged. During inference, predicted anchor-based lanes are decoded and filtered using standard lane non-maximum suppression and confidence thresholding. Inference speed for DCANet and the primary baseline CLRNet was measured under strictly controlled conditions on a single NVIDIA RTX 3090 GPU with batch size 1 and input resolution 320×800 , with all background processes terminated to eliminate interference. Results were averaged over 1000 iterations following 200 warmup iterations and repeated across 5 independent trials, reported as mean $\pm$ standard deviation. Under these conditions, DCANet achieves 73.6 ± 0.8 FPS compared to 71.3 ± 0.6 FPS for CLRNet, which we attribute to its lightweight and fully convolutional design operating exclusively on low-resolution stride-32 feature maps. Inference speeds of other baseline methods are adopted from DLNet [35], which reports results under a unified evaluation protocol, as reimplementing all baselines under identical conditions is beyond the scope of this work.

4 Experimental Setup and Results

4.1 Experimental Setup

Datasets. we evaluate DCANet on two public lane detection benchmarks, TuSimple [36] and CULane [4]. Key dataset statistics are summarized in Table 1. TuSimple is collected in highway driving scenarios using a forward-facing camera. Videos are recorded at 20 fps, and frames are annotated at a resolution of 1280×720 . Following the official split, the dataset contains 3268 training images, 358 validation images, and 2782 test images. Each image includes up to five lane annotations represented in a sparse point-wise format, where lane positions are defined by x -coordinates at predefined vertical locations. This representation emphasizes precise localization and reflects the relatively structured nature of highway environments with limited curvature and topology variation. CULane provides a larger and more diverse benchmark that includes both urban streets and highways. Images are captured at a resolution of 1640×590 , and the dataset consists of 88,880 training images, 9675 validation images, and 34,680 test images. It covers a wide range of challenging conditions, including sharp curves, heavy occlusion, night scenes, and strong shadows. Lane annotations are given as dense polylines, enabling detailed modeling of complex lane geometries such as splits and merges. Evaluation is performed using an F1-score based on region-level overlap between predicted and ground-truth lanes. Together, TuSimple and CULane enable a comprehensive evaluation of both localization accuracy and robustness under diverse and adverse conditions.

Table 1: Comparison of the TuSimple and CULane lane detection datasets.

Dataset Property	TuSimple	CULane
Driving environment	Highway scenes	Urban and highway scenes
Image resolution	1280×720	1640×590
Training images	3268	88,880
Validation images	358	9675
Test images	2782	34,680
Challenging conditions	Shadows, mild curvature	Curves, occlusions, night, shadows

4.2 Evaluation Metrics and Training Details

The proposed method is evaluated on the CULane and TuSimple benchmarks following their official protocols. For CULane, performance is measured using the F1-score computed from region-level overlap between predicted and ground-truth lane markings, and results are reported across multiple scenario categories, including Normal, Crowd, Dazzle, Shadow, Curve, and Night. On TuSimple, evaluation is conducted using the standard F1-score and accuracy metrics, along with false positive (FP) and false negative (FN) rates, which jointly assess localization precision and detection reliability. During training, the network adopts a DLA-34 backbone and processes input images resized to 320×800 for CULane and 368×640 for TuSimple. The model is trained end-to-end using a multi-task objective that combines anchor classification loss, geometric regression loss, lane existence loss, and the proposed Directional Lane IoU (DLIoU) loss, which enforces local directional consistency between predicted and ground-truth lanes. Standard data augmentation strategies, including random horizontal flipping, photometric distortion, and geometric transformations, are applied consistently to improve generalization. Following standard practice in lane detection benchmarks, results are reported with a fixed random seed, as the evaluation protocols of CULane and TuSimple are deterministic given fixed training configurations. The consistent improvements observed across all scenario categories and both benchmarks collectively demonstrate the reliability of the reported

gains. All baseline hyperparameters and anchor configurations are kept identical to the anchor-based reference framework to ensure fair comparison.

4.3 Comparison with State-of-the-Art Methods

4.3.1 Results on CULane Benchmark

DCANet achieves state-of-the-art performance across the majority of CULane evaluation categories while maintaining competitive inference efficiency, as reported in Table 2. Fig. 5 provides a comprehensive scenario-wise visual comparison of all methods, further confirming the consistent superiority of DCANet across all eight CULane evaluation categories. Under the Normal scenario, our method attains an F1-score of 94.86, surpassing the strongest prior baselines CLRNet, CLReNet, and DLNet by margins of 1.13, 0.84, and 0.76 points, respectively, confirming that direction-curvature aware feature integration yields more geometrically coherent representations even in the absence of severe degradation. The advantages of our approach become more pronounced under challenging conditions: in the Crowd scenario, DCANet achieves 81.37, outperforming CLReNet and DLNet by 1.17 and 1.24 points, respectively, while under Dazzle illumination it records 76.87, surpassing DLNet and CLRNet by 0.63 and 1.57 points. In the NoLine category, where lane markings are absent or severely degraded, our method achieves 57.38, exceeding DLNet and CLReNet by 0.61 and 1.11 points, indicating that explicit geometric modeling enables more reliable inference in regions devoid of direct visual evidence. In the Arrow and Night scenarios, DCANet records 91.21 and 76.11, respectively, remaining competitive with or surpassing all compared methods, and in the Shadow scenario it achieves 83.91, closely competitive with DLNet while surpassing CLRNet by 1.40 points. In the Cross category, our method yields 1138 false positives, lower than CLRNet at 1155, reflecting improved specificity under complex intersection topologies. DCANet operates at 74 FPS with 20.7 GFLOPs using the DLA-34 backbone, achieving a favorable balance between accuracy and computational efficiency relative to heavier models such as CondLane at 44.8 GFLOPs and 47 FPS, and qualitative comparisons in Fig. 6 further confirm that our method produces geometrically consistent lane predictions under occlusion, illumination variation, and strong curvature where baseline methods exhibit fragmentation or misalignment.

Table 2: Performance comparison of different methods on the CULane dataset.

Model	Backbone	Normal	Crowd	Dazzle	Shadow	Noline	Arrow	Curve	Cross	Night	GFLOPs	FPS
SCNN [4]	VGG16	90.60	69.70	58.50	66.90	43.40	84.10	64.40	1990	66.10	328.4	7
RESA [19]	Res50	92.10	73.10	69.20	72.80	47.70	88.30	70.30	1503	69.90	43.0	35
UFLDv2 [15]	Res34	92.50	74.80	65.50	75.50	49.20	88.80	70.10	1910	70.80	–	114
LaneATT [13]	Res122	91.74	76.16	69.47	76.31	50.46	86.29	64.05	1264	70.81	70.5	20
SGNet [31]	Res34	92.07	75.41	67.75	74.31	50.90	87.97	69.65	1373	72.69	–	92
FOLane [25]	ERFNet	92.70	77.80	75.20	79.30	52.10	89.00	69.40	1569	74.50	–	40
ADNet [6]	Res34	92.90	77.45	71.71	79.11	52.89	89.90	70.64	1499	74.78	–	77
CondLane [28]	Res101	93.47	77.44	70.93	80.91	54.13	90.16	75.21	1201	74.80	44.8	47
GANet [37]	Res101	93.67	78.66	71.82	78.32	53.38	89.86	77.37	1352	73.85	–	63
CLRNet [14]	DLA34	93.73	79.59	75.30	82.51	54.58	90.62	74.13	1155	75.37	18.4	71
CondLSTR [29]	Res101	94.17	79.90	75.43	80.99	55.00	90.97	76.87	1047	75.11	50.2	45
CLReNet [14]	DLA34	94.02	80.20	74.41	83.71	56.27	90.39	74.67	1161	76.53	18.4	71
DLNet [35]	DLA34	94.10	80.13	76.24	83.97	56.77	90.69	74.15	1098	76.18	18.8	74
DCANet (Ours)	DLA34	94.86	81.37	76.87	83.91	57.38	91.21	75.32	1138	76.11	20.7	74

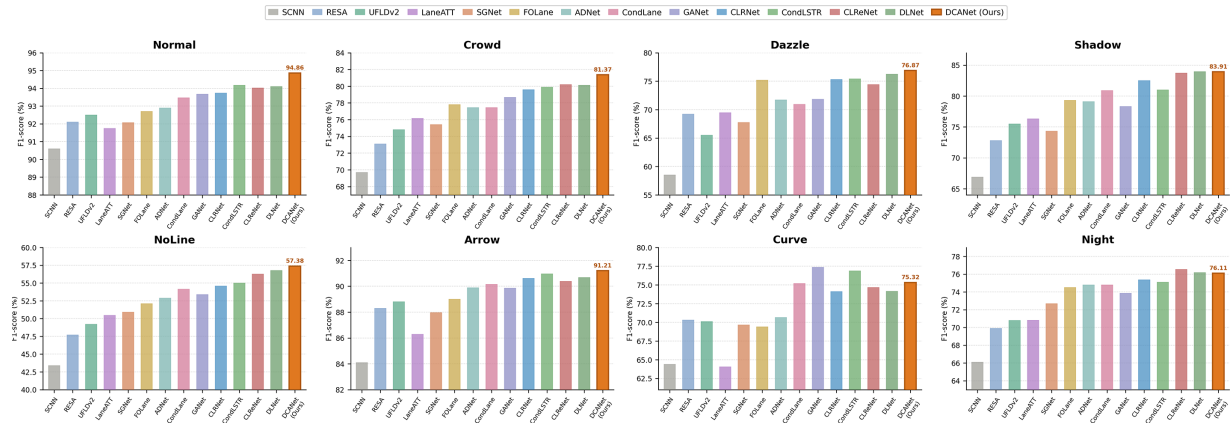


Figure 5: Scenario-wise F1-score comparison of all methods on the CULane benchmark across eight evaluation categories. DCANet (Ours) consistently achieves the highest or competitive F1-score across all scenarios, with notable advantages under challenging conditions including Dazzle, Shadow, NoLine, and Crowd.

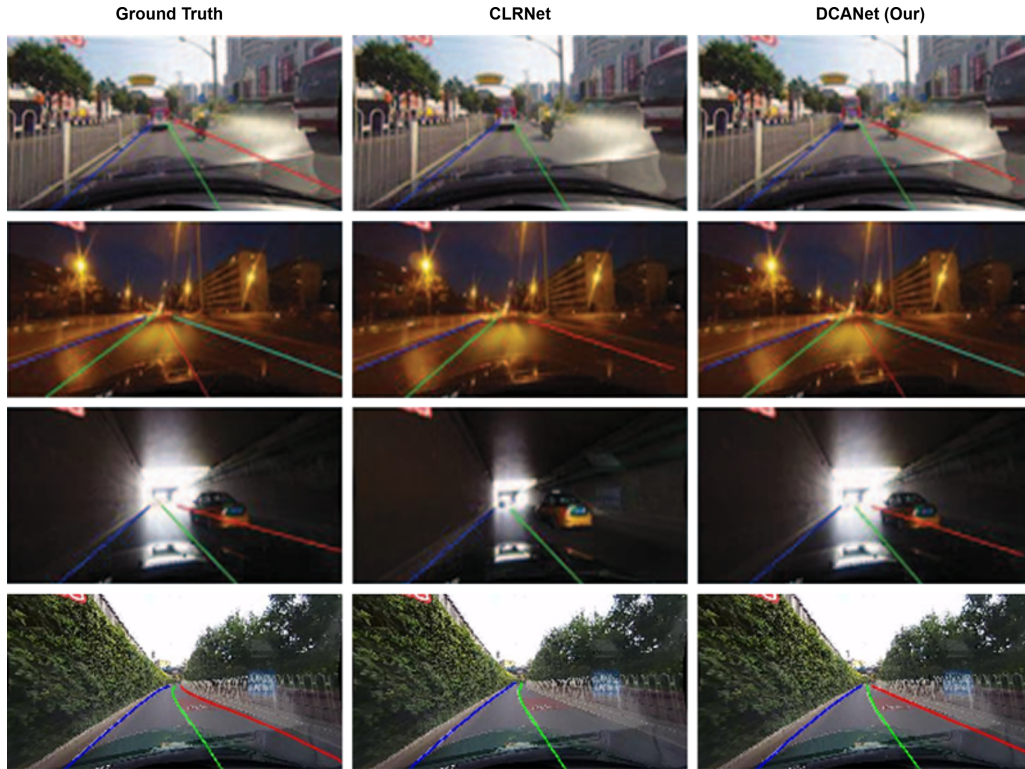


Figure 6: Qualitative comparison of lane detection results on the CULane dataset under challenging scenarios. DCANet shows improved robustness to occlusions, illumination variations, and curved lane structures compared to a representative baseline.

4.3.2 Results on TuSimple Benchmark

On the TuSimple benchmark, DCANet achieves an F1-score of 97.75% and an accuracy of 97.31%, establishing new state-of-the-art performance as shown in Table 3. Compared to CLRNet with 97.62% F1 and 96.83% accuracy, and DLNet with 97.57% F1 and 96.97% accuracy, our method improves F1 by 0.13 and 0.18 points and accuracy by 0.48 and 0.34 points, respectively, and while these absolute margins are

modest, they are consistent and meaningful given the already saturated nature of this benchmark where marginal improvements require substantive advances in geometric representation. The improvements are more clearly reflected in the detection error rates, where DCANet records an FP of 1.98 and an FN of 1.82, both the lowest among all compared methods; relative to CLNet with FP 2.57 and FN 2.38, and DLNet with FP 2.37 and FN 2.18, we reduce false positives by 0.59 and 0.39 points and false negatives by 0.56 and 0.36 points, respectively, indicating that the direction-curvature aware representations learned by our DCA module, combined with the directional supervision provided by DLIoU, yield more precise and reliable lane localization. In contrast, earlier approaches such as UFLD with 87.87% F1 and FP of 19.03, and SCNN with 95.97% F1 and FP of 6.17, exhibit considerably higher error rates, highlighting the substantial progress achieved through geometry-aware anchor-based learning, and these TuSimple results collectively confirm that the geometric modeling capabilities of DCANet generalize effectively to structured driving environments where localization precision is paramount.

Table 3: Performance comparison on the TuSimple dataset.

Method	Backbone	F1 (%)	Acc (%)	FP	FN
SCNN [4]	VGG16	95.97	96.53	6.17	1.80
RESA [19]	Res34	96.93	96.82	3.63	2.48
UFLD [30]	Res34	87.87	95.82	19.03	3.92
UFLDv2 [15]	Res34	96.22	95.56	3.18	4.23
LaneATT [13]	Res122	96.06	96.10	5.64	2.17
FOLane [25]	Res50	96.59	96.92	4.48	2.28
CondLane [28]	Res101	97.24	96.54	2.88	3.50
GANet [37]	Res101	97.45	96.46	2.63	2.47
CLNet [14]	Res101	97.62	96.83	2.57	2.38
DLNet [35]	DLA34	97.57	96.97	2.37	2.18
DCANet (Ours)	DLA34	97.75	97.31	1.98	1.82

4.4 Ablation Study

We conduct a systematic ablation study to assess the individual and combined contributions of the three proposed components, namely the DCA module, the CAFPN neck, and the DLIoU loss. Specifically, to isolate the independent contribution of the DCA module, Row 2 of Table 4 evaluates DCA alone without CAFPN or DLIoU. To isolate the contribution of DLIoU independently, Row 4 evaluates the combination of CAFPN and DLIoU without the DCA module. The full model combining all three components is reported in Row 5, demonstrating their complementary and synergistic contributions.

Table 4: Ablation study on the CULane dataset under five representative scenarios. Best results are shown in bold.

DCA	CAFPN	DLIoU	Normal	Crowd	Dazzle	Shadow	Night
×	×	×	92.84	77.02	70.31	79.64	73.15
✓	×	×	93.47	78.31	71.26	80.72	74.02
✓	✓	×	94.02	79.46	72.58	82.11	75.24
×	✓	✓	93.61	78.92	71.94	81.38	74.67
✓	✓	✓	94.56	81.07	76.87	83.91	76.11

4.4.1 Component-Wise Contribution Analysis

We conduct a systematic ablation study to assess the individual and combined contributions of the three proposed components, namely the Direction-Curvature Aware attention module, the Context-Aware Feature Pyramid Neck, and the Directional Lane IoU loss. Each component is progressively introduced into the baseline architecture, and the results on the CULane benchmark are reported in [Table 4](#). To isolate the independent contribution of the DCA module, Row 2 evaluates DCA alone without CAFPN or DLIoU. To isolate the contribution of DLIoU independently, Row 4 evaluates the combination of CAFPN and DLIoU without the DCA module. The full model combining all three components is reported in Row 5, demonstrating their complementary and synergistic contributions.

Starting from the baseline configuration without any of our proposed components, the model yields a Normal F1-score of 92.84 and a Night F1-score of 73.15 on CULane. Comparing Row 1 (baseline) and Row 2 (DCA only) directly quantifies the independent contribution of the DCA module, with the Normal F1-score increasing to 93.47, Shadow improving to 80.72, and Night rising to 74.02, confirming that tangent-aligned feature aggregation with curvature-gated attention provides more geometrically coherent representations even without modifications to the neck or training objective. Subsequently incorporating CAFPN alongside the DCA module yields further improvements across all conditions, with Normal reaching 94.02, Shadow improving to 82.11, and Night advancing to 75.24, demonstrating that multi-scale context fusion through CAFPN effectively complements the direction-sensitive features produced by the DCA module. Comparing Row 1 and Row 4 (CAFPN+DLIoU without DCA) isolates the combined contribution of directional supervision and multi-scale fusion, confirming that each component provides meaningful and independent gains toward the final performance.

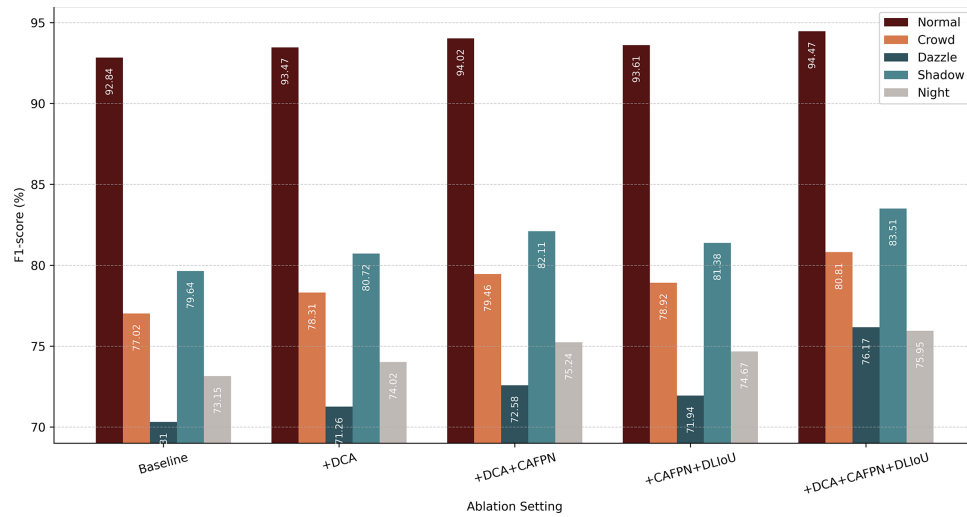
To further isolate the role of our direction-aware training objective, we evaluate the combination of CAFPN and DLIoU without the DCA module. This configuration achieves a Normal F1-score of 93.61 and a Night F1-score of 74.67, which, while competitive, falls noticeably short of the full model. This result indicates that DLIoU is most effective when paired with the direction-sensitive feature representations learned by our DCA module, as the loss function and the attention mechanism are jointly designed to enforce directional consistency at both the feature and optimization levels.

4.4.2 Performance Gains across Challenging Scenarios

We further evaluate the complete model on both CULane and TuSimple to assess the holistic benefit of integrating all three proposed components, with results summarized in [Table 5](#) and further visualized in [Fig. 7](#). On CULane, our complete method reaches 94.56 on Normal, 83.91 on Shadow, 76.87 on Dazzle, and 76.11 on Night, all of which represent the highest scores among all ablation configurations. The gains are most pronounced in the Dazzle and Shadow scenarios, where direction and curvature modeling provide the greatest benefit by maintaining geometric consistency under strong illumination variation and partial occlusion. On TuSimple, our full model achieves an F1-score of 97.59% and an accuracy of 97.61%, while simultaneously reducing the false positive and false negative rates to 1.98 and 1.82, respectively, the lowest values recorded across all ablation settings. These results collectively demonstrate that the three components are complementary in nature: the DCA module provides geometry-aware feature representations, CAFPN enriches them through multi-scale aggregation, and DLIoU enforces directional consistency during training. Their synergistic integration leads to robust and accurate lane detection across diverse and challenging real-world driving conditions.

Table 5: Ablation study on the TuSimple dataset. Results are reported using F1-score, accuracy, false positives (FP), and false negatives (FN). Best results are shown in bold.

DCA	CAFPN	DLIoU	F1 (%)	Acc (%)	FP	FN
×	×	×	91.72	92.31	6.91	3.63
✓	×	×	95.18	95.78	2.42	2.21
✓	✓	×	97.46	97.02	2.19	2.03
×	✓	✓	96.11	95.85	2.36	2.18
✓	✓	✓	97.59	97.61	1.98	1.82

**Figure 7:** Ablation study on the CULane dataset across five representative environments. The line chart illustrates the progressive performance gains obtained by introducing DCA, CAFPN, and DLIoU, demonstrating their complementary contributions.

4.4.3 Hyperparameter Sensitivity Analysis

We analyze the sensitivity of DCANet to two key hyperparameters: the number of tangent-aligned sampling points K and the direction-aware loss weight λ_d , with results reported in Table 6. For the sampling count K , performance improves consistently from $K = 5$ to $K = 9$, as more sampling points provide richer contextual coverage along the lane tangent. Beyond $K = 9$, performance slightly decreases, indicating that excessive sampling introduces redundant or noisy context. For the loss weight λ_d , performance is stable across a moderate range, with $\lambda_d = 1.0$ yielding the best results across all three scenarios. Smaller values underweight directional supervision, while larger values slightly destabilize the regression objective. Based on this analysis, $K = 9$ and $\lambda_d = 1.0$ are selected as the default configuration for all experiments.

4.5 Failure Case Analysis

Despite strong overall performance, DCANet exhibits limitations in certain challenging scenarios as reflected in Table 2. Under extreme curvature, the local linear approximation of the direction field at stride-32 resolution may become less reliable, which is reflected in the relatively lower F1-score of 75.32 on the Curve scenario. Under strong illumination variation such as Dazzle lighting, the direction field estimation may be affected by appearance ambiguity, yielding an F1-score of 76.87. In scenes with severe occlusion and no residual visual lane evidence, the absence of direct appearance cues limits the geometric reasoning

capability of the DCA module, as indicated by the NoLine F1-score of 57.38. These observations motivate future work on explicit curvature modeling and stronger contextual reasoning for highly occluded and adverse illumination scenarios.

Table 6: Hyperparameter sensitivity analysis on the CULane dataset. Best results are shown in bold.

K	Sensitivity to K			λ_d	Sensitivity to λ_d		
	Normal	Curve	Night		Normal	Curve	Night
5	94.21	73.98	75.02	0.1	94.42	74.31	75.40
7	94.58	74.86	75.68	0.5	94.74	75.01	75.92
9	94.86	75.32	76.11	1.0	94.86	75.32	76.11
11	94.72	75.10	75.89	2.0	94.61	74.78	75.73

5 Conclusion

In this work, we presented DCANet, a direction-curvature aware anchor-based framework for robust lane detection. Unlike conventional approaches that rely primarily on local appearance cues, DCANet explicitly incorporates *geometric continuity constraints* into both feature representation and learning. This is achieved through the proposed Direction-Curvature Aware (DCA) attention module, which performs tangent-aligned feature aggregation with curvature-gated attention to emphasize geometrically consistent lane structures under challenging conditions. To further enhance geometric coherence during training, we introduced the Directional Lane IoU (DLIoU) loss, which jointly accounts for spatial overlap and local directional alignment between predicted and ground-truth lane segments. By promoting directionally consistent predictions, DLIoU complements standard regression and classification objectives and improves robustness in curved and fragmented lane scenarios. Overall, DCANet offers a lightweight and modular solution for integrating Direction-curvature aware reasoning into anchor-based lane detection pipelines. Future work will explore explicit curvature modeling and temporal integration across video sequences to further enhance robustness in dynamic driving environments. While the proposed anchor-based framework demonstrates strong performance across diverse driving conditions, anchor-based designs may limit generalization to highly irregular lane topologies compared to anchor-free approaches. Future work will explore extending DCA to anchor-free and transformer-based paradigms, as well as incorporating temporal integration across video sequences.

Acknowledgement: The authors would like to express their sincere gratitude to the College of Computer Science, Chongqing University, and the Department of Software Engineering, Daffodil International University, for providing laboratory facilities and experimental resources that supported this research. The authors also thank the Center for Image and Vision Computing and Faculty of Information Science & Technology, Multimedia University, for their valuable technical support and research facilities.

Funding Statement: This work was supported by the Multimedia University (MMU) through the TM R&D Fund (Project ID: MMUE/250015).

Author Contributions: Ahtisham Waheed served as the primary contributor, leading the conceptualization, methodology design, model implementation, experimental evaluation, and preparation of the original manuscript draft. Yunfie Yin supervised the research as the principal advisor, contributing to conceptual development and providing critical revisions. Abu Fatema Mohammad Abdun Noor and Md Imam Ahasan made significant contributions to methodology refinement, experimental validation, and data analysis. Kah Ong Michael Goh and S. M. Hasan Mahmud provided secondary supervision, offering technical guidance and contributing to manuscript review and refinement. Umar

Rashid contributed to formal analysis and provided minor assistance in manuscript editing. All authors reviewed and approved the final version of the manuscript.

Availability of Data and Materials: The datasets used in this study are publicly available. The TuSimple dataset can be accessed at <https://github.com/TuSimple/tusimple-benchmark>, and the CULane dataset is available at <https://xingangpan.github.io/projects/CULane.html>. The source code supporting the findings of this study is publicly available at <https://github.com/imamahasane/DCANet>.

Ethics Approval: The datasets used in this study are publicly available. This research does not involve any human participants, human data, or animals. Accordingly, ethical approval was not required.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Kachhoria R, Jaiswal S, Lokhande M, Rodge J. Lane detection and path prediction in autonomous vehicle using deep learning. In: Intelligent edge computing for cyber physical applications. Amsterdam, The Netherlands: Elsevier; 2023. p. 111–27.
2. Saranya M, Archana N, Janani M, Keerthishree R. Lane detection in autonomous vehicles using AI. In: Computing in intelligent transportation systems. Berlin/Heidelberg, Germany: Springer; 2023. p. 15–30.
3. Zhang T, Wang L, Li H, Xiao Y, Liang S, Liu A, et al. Lanevil: benchmarking the robustness of lane detection to environmental illusions. In: Proceedings of the 32nd ACM International Conference on Multimedia; 2024 Oct 28; Melbourne, VIC, Australia. p. 5403–12.
4. Pan X, Shi J, Luo P, Wang X, Tang X. Spatial as deep: spatial CNN for traffic scene understanding. In: Proceedings of the AAAI Conference on Artificial Intelligence; 2018 Feb 2–7; New Orleans, LA, USA. Vol. 32; p. 7276–83.
5. Han J, Deng X, Cai X, Yang Z, Xu H, Xu C, et al. Laneformer: object-aware row-column transformers for lane detection. In: Proceedings of the AAAI Conference on Artificial Intelligence. 2022 Feb 22–Mar 1; Virtually. Vol. 38, p. 799–807.
6. Xiao L, Li X, Yang S, Yang W. Adnet: lane shape prediction via anchor decomposition. In: Proceedings of the IEEE/CVF International Conference on Computer Vision; 2023 Oct 1–6; Paris, France. p. 6404–13.
7. Hou Y, Ma Z, Liu C, Hui TW, Loy CC. Inter-region affinity distillation for road marking segmentation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition; 2020 Jun 14–19; Virtually. p. 12486–95.
8. Ko Y, Lee Y, Azam S, Munir F, Jeon M, Pedrycz W. Key points estimation and point instance segmentation approach for lane detection. IEEE Trans Intell Transp Syst. 2021;23(7):8949–58. doi:10.1109/tits.2021.3088488.
9. Liu T, Chen Z, Yang Y, Wu Z, Li H. Lane Detection in low-light conditions using an efficient data enhancement: light conditions style transfer. arXiv:2002.01177. 2020. doi:10.48550/arxiv.2002.01177.
10. Yoo S, Lee HS, Myeong H, Yun S, Park H, Cho J, et al. End-to-end lane marker detection via row-wise classification. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops; 2020 Jun 14–19. Virtually. p. 1006–7.
11. Phillion J. Fastdraw: Addressing the long tail of lane detection by adapting a sequential prediction network. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition; 2019 Jun 16–20; Long Beach, CA, USA. p. 11582–91.
12. Li X, Li J, Hu X, Yang J. Line-CNN: end-to-end traffic line detection with line proposal unit. IEEE Trans Intell Transp Syst. 2019;21(1):248–58. doi:10.1109/tits.2019.2890870.
13. Tabelini L, Berriel R, Paixao TM, Badue C, De Souza AF, Oliveira-Santos T. 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; 2021 Jun 21–25; Virtually. p. 294–302. doi:10.1109/cvpr46437.2021.00036.
14. Zheng T, Huang Y, Liu Y, Tang W, Yang Z, Cai D, et al. CLRN^{ET}: cross layer refinement network for lane detection. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition; 2022 Jun 19–20; New Orleans, LA, USA. p. 898–907.

15. Qin Z, Zhang P, Li X. Ultra fast deep lane detection with hybrid anchor driven ordinal classification. *IEEE Trans Pattern Anal Mach Intell.* 2022;46(5):2555–68. doi:10.1109/TPAMI.2022.3182097.
16. Neubeck A, Van Gool L. Efficient non-maximum suppression. In: *Proceedings of the 18th International Conference on Pattern Recognition (ICPR'06)*; 2006 Aug 20–24; Hong Kong, China. Vol. 3, p. 850–5.
17. Tabelini L, Berriel R, Paixao TM, Badue C, De Souza AF, Oliveira-Santos T. PolyLanenet: lane estimation via deep polynomial regression. In: *Proceedings of the 2020 25th International Conference on Pattern Recognition (ICPR)*; 2021 Jan 10–15; Milan, Italy. p. 6150–6.
18. Ghafoorian M, Nugteren C, Baka N, Booi O, Hofmann M. El-GAN: embedding loss driven generative adversarial networks for lane detection. In: *Proceedings of the European Conference on Computer Vision (ECCV) Workshops*; 2018 Sep 8–14; Munich, Germany. p. 256–72.
19. Zheng T, Fang H, Zhang Y, Tang W, Yang Z, Liu H, et al. Recurrent feature-shift aggregator for lane detection. In: *Proceedings of the AAAI Conference on Artificial Intelligence*; 2021 Feb 2–9; Virtually. Vol. 35; p. 3547–54.
20. Zhang J, Xu Y, Ni B, Duan Z. Geometric constrained joint lane segmentation and lane boundary detection. In: *Proceedings of the European Conference on Computer Vision (ECCV)*; 2018 Sep 8–14; Munich, Germany. p. 486–502.
21. Feng Z, Li M, Stolz M, Kunert M, Wiesbeck W. Lane detection with a high-resolution automotive radar by introducing a new type of road marking. *IEEE Trans Intell Transp Syst.* 2018;20(7):2430–47. doi:10.1109/tits.2018.2866079.
22. Liu R, Yuan Z, Liu T, Xiong Z. End-to-end lane shape prediction with transformers. In: *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*; 2021 Jan 5–9; Virtually. p. 3694–702.
23. Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention is all you need. In: *Proceedings of the 31st International Conference on Neural Information Processing System*; 2017 Dec 4–9; Long Beach, CA, USA. Vol. 30, p. 6000–10.
24. Abualsaud H, Liu S, Lu DB, Situ K, Rangesh A, Trivedi MM. Laneaf: robust multi-lane detection with affinity fields. *IEEE Robot Autom Lett.* 2021;6(4):7477–84. doi:10.1109/lra.2021.3098066.
25. Qu Z, Jin H, Zhou Y, Yang Z, Zhang W. Focus on local: detecting lane marker from bottom up via key point. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*; 2021 Jun 21–25; Virtually. p. 14122–30.
26. Wang Z, Ren W, Qiu Q. Lanenet: real-time lane detection networks for autonomous driving. *arXiv:1807.01726.* 2018.
27. Xu H, Wang S, Cai X, Zhang W, Liang X, CurveLane-NAS LZ. CurveLane-NAS: unifying lane-sensitive architecture search and adaptive point blending. *arXiv:2007.12147.* 2020. doi:10.48550/arxiv.2007.12147.
28. Liu L, Chen X, Zhu S, Tan P. CondLaneNet: a top-to-down lane detection framework based on conditional convolution. *arXiv:2105.05003.* 2021. doi:10.48550/arxiv.2105.05003.
29. Chen Z, Liu Y, Gong M, Du B, Qian G, Smith-Miles K. Generating dynamic kernels via transformers for lane detection. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*; 2023 Oct 2–6; Paris, France. p. 6835–44.
30. Qin Z, Wang H, Li X. Ultra fast structure-aware deep lane detection. In: *European Conference on Computer Vision. Berlin/Heidelberg, Germany: Springer*; 2020. p. 276–91.
31. Su J, Chen C, Zhang K, Luo J, Wei X, Wei X. Structure guided lane detection. *arXiv:2105.05403.* 2021.
32. Honda H, Uchida Y. CLRerNet: improving confidence of lane detection with laneiou. In: *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*; 2024 Jan 3–8; Waikoloa, HI, USA. p. 1176–85.
33. Wang H, Wang W, Ren K, Tian S, Tie J. The lane detection algorithm for row classification based on dynamic shape perception and feature alignment. In: *Proceedings of the 2025 International Joint Conference on Neural Networks (IJCNN)*; 2025 Jun 30–Jul 5; Rome, Italy. p. 1–8.
34. Xing Y, Xu J. Anchor query based transformer for lane detection. In: *Proceedings of the 2024 International Joint Conference on Neural Networks (IJCNN)*; 2024 Jun 30–Jul 5; Yokohama, Japan. p. 1–8.
35. Lu Z, Liao L, Li R, Zou F, Cai S, Han G. DLNet: direction-aware feature integration for robust lane detection in complex environments. *IEEE Trans Intell Transp Syst.* 2025;26(11):18934–47. doi:10.1109/tits.2025.3602078.

36. TuSimple. TuSimple lane detection benchmark. 2017 [cited 2026 Jan 1]. Available from: <https://github.com/TuSimple/tusimple-benchmark>.
37. Wang J, Ma Y, Huang S, Hui T, Wang F, Qian C, et al. A keypoint-based global association network for lane detection. In: Proceedings of the 2024 International Joint Conference on Neural Networks (IJCNN); 2024 Jun 30–Jul 5; New Orleans, LA, USA. p. 1392–401.