



International Journal of Research Publication and Reviews

Journal homepage: www.ijrpr.com ISSN 2582-7421

Lane Detection Using Hybrid Modeling

Jyoti¹, Amandeep², Isha Bisht³, Deepak Kumar⁴, Satbir⁵

Email:-Jyotipectan@Gmail.Com

¹M.Sc Cs (Ai &Ds),Guru Jambheshwar University Of Science And Technology, Hisar

²Assistant Professor(Ai &Ds),Guru Jambheshwar University Of Science And Technology, Hisar

^{3,4,5}M.Sc Cs (Ai &Ds),Guru Jambheshwar University Of Science And Technology, Hisar

DOI : <https://doi.org/10.5281/zenodo.15735267>

ABSTRACT—

The paper introduces a new lane detection model for autonomous cars and advanced driver assistance systems (ADAS) that improves on traditional methods struggling with tough conditions like shadows, worn-out lane markings, or occlusions. Unlike single-frame approaches, this model combines classic vision techniques with modern deep learning for better accuracy and real-time performance. It uses a Deep Convolutional Neural Network (DCNN) with an encoder-decoder setup to analyze spatial details in each frame and a Deep Recurrent Neural Network (DRNN) with Convolutional Long Short-Term Memory (ConvLSTM) units to leverage temporal connections across frames. A hybrid attention module focuses on lane-specific features, and a fusion reasoning system blends rule-based and deep learning outputs for robust results. The model was tested on the TuSimple dataset and two custom datasets (urban and rural roads), targeting over 60 FPS, an IoU above 0.3, and over 95% accuracy on edge devices. This approach enhances scalability, accuracy, and speed for safer autonomous driving by moving beyond single-frame limitations and optimizing computation.

Keywords—Convolutional neural network, LSTM, lane detection, semantic segmentation, autonomous driving.

1. INTRODUCTION

The project aims to develop an advanced lane detection system for autonomous vehicles (ADAS) by integrating traditional computer vision techniques with deep learning to achieve robust, real-time performance. Traditional methods, such as edge detection and Hough transforms, are computationally efficient but struggle in challenging conditions like heavy shadows, faded road markings, or occlusions [2, 3, 9]. Deep learning approaches, are particularly convolutional neural networks (CNNs), excel in semantic understanding but often require significant computational resources, limiting their use on resource-constrained edge devices [6, 10]. The proposed hybrid model combines a (DCNN) with an encoder-decoder architecture to extract spatial features from individual frames [4] and a Deep Recurrent Neural Network (DRNN) with Long Short-Term Memory (LSTM) units to process multiple frames as a time-series, leveraging temporal continuity for improved lane detection [4, 9]. The model incorporates preprocessing techniques like Canny edge detection and inverse perspective mapping, alongside advanced deep learning features such as hybrid attention mechanisms and online re-parameterization, to balance accuracy, robustness, and computational efficiency [11, 12]. This hybrid approach enhances performance across diverse scenarios, including highways, urban streets, and unstructured

rural roads [5, 8, 13]. To ensure comprehensive evaluation, the project introduces two new datasets—one capturing diverse driving conditions and another focusing on rural roads—complementing existing datasets like TuSimple and CULane [6, 15]. The system is optimized for real-time operation on edge devices, with performance measured using metrics such as frames per second (FPS), precision, and intersection over union (IoU) [14, 15].

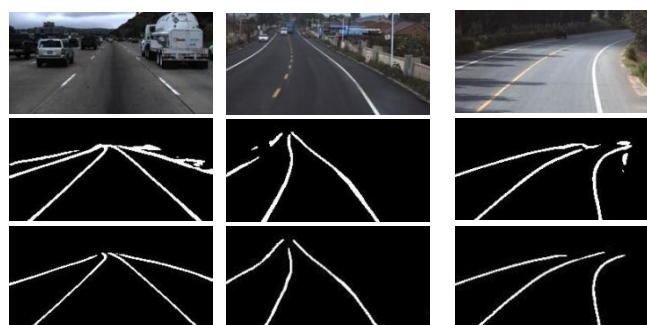


Figure 1.1. Three sample photos of various road scenarios

The goal is to create a scalable, accurate, and adaptable lane detection system that supports safe navigation, lane departure warnings, and path planning in complex driving environments [1, 7]. Lane detection is vital for autonomous driving and (ADAS), enabling vehicles to track road boundaries, support lane departure warnings, and plan paths [1]. However, traditional methods, such as edge detection and Hough transforms, struggle in complex scenarios like heavy shadows, faded markings, or occlusions, while deep learning models, though powerful for semantic understanding, are computationally intensive, limiting real-time performance on edge devices [2]. This project proposes a hybrid lane detection model that integrates the efficiency of traditional vision techniques with the robustness of deep learning to achieve accurate, real-time performance, even in challenging conditions [3]. To address these challenges, the model integrates traditional techniques—such as Canny edge detection, inverse perspective mapping, Hough transforms—for early stage lane candidate finding [4] with a deep learning pipeline including both a CNN and (DRNN) [5]. The CNN, constructed following the encoder-decoder paradigm with a re-parameterized ResNet-18, extracts spatial features (e.g., edges, lane shapes) from the individual frame [6], while the DRNN, which uses (LSTM) units, processes them over a sequence of adjacent frames to capture

temporal continuity and improve the detection in dynamic scenes such as occlusive frames or degraded markings [7]. A hybrid attention-method that integrates positional and channel attention enhances the focus on lane-specific features [8] and a fusion approach integrates classical and DL-fused outputs to produce more robust lane predictions [9]. The model produces binary map and instance-pertinent lane points, which are further refined by post-processing means such as polynomial fitting and temporal smoothing [10]. The encoder in CNN [11] which is built on ResNet, quickly makes features become hierarchical thereby incising the feature depth dimension as well as decreasing the spatial size and can be optimized for edge device with no need for offline training and fine-tuning. The decoder consists of a segmentation branch to distinguish lane from non-lane pixels and a row-anchor classification branch to separate lane instances [12]. A loss integrating segmentation instance clustering and smooth curve regularization. We propose a loss that combines binary cross-entropy for segmentation, a discriminative loss for instance clustering, and an optional polynomial fitting loss for smooth lane curves [13]. The DRNN's LSTM processes time-series feature maps, making the model robust to challenging conditions [14].

Lane detection is crucial for autonomous vehicles and ADAS, ensuring vehicles stay within their lanes and avoid collisions [1]. Traditional methods, including geometric modeling [2, 3], energy-based techniques [4, 5], and single-frame deep learning [6–9], often struggle in tough conditions like heavy shadows, faded road markings, or occlusions, as single images lack sufficient context. Since road lanes are continuous and driving scene frames are highly correlated, using multiple frames can enhance detection accuracy by providing contextual clues from prior frames [1]. The proposed approach introduces a hybrid deep neural network combining (DCNNs) (DRNNs) to process multiple consecutive frames [1]. DCNNs extract compact feature maps from high-resolution images (e.g., 800x600 pixels), which are then fed into a DRNN with (LSTM) units to analyze time-series data across frames [7]. The network operates as a DCNN with an encoder-decoder architecture for semantic segmentation, ensuring output maps match the input image size [10]. The encoder generates feature maps, the LSTM integrates temporal information, and the decoder produces accurate lane predictions [1]. The project highlights three key contributions:

Multi-frame Lane Detection: By analyzing a sequence of frames, this method outperforms single-frame approaches, especially in challenging scenarios like shadows or occlusions [1, 9].

Hybrid Network Architecture: Seamlessly integrates DCNNs for spatial feature extraction and LSTM-based DRNNs for temporal processing, enhancing lane prediction accuracy [1].

New Datasets: Includes two new datasets—one with diverse driving scenarios and another focused on rural roads—alongside an expanded TuSimple dataset, enabling robust testing [1].

2. RELATED WORK

Convolutional-Based Approaches for Lane Detection:-

CNNs have revolutionized image segmentation and lane detection by providing effective feature extraction and pattern recognition for autonomous driving and advanced driver assistance systems (ADAS) [1]. CNNs have been successful in detecting road markings (solid vs dashed lines) and directly processing images to produce the lane boundary, in challenging conditions, such as shadow or occlusion, in real-time.

An alternative new approach from Getahun et al. (2021) employs a CNN with sliding window technique to localize lane markings using small image sub-regions [2]. The architecture of the model includes three convolution layers, one max-pooling layer, three layers of fully connected layer with ReLU activation and a Softmax sigmoid for binary classification between the lane and the background [2]. Post-processing filters low-confidence patches, computes average lane points, and clusters them using Euclidean distance and angle to group points into lane boundaries, achieving robust performance at 28 FPS in complex scenarios [2].

Another approach by Tabelini et al. (2022) leverages temporal information from sequential frames, enhancing the LaneATT model [3, 4]. LaneATT includes feature extraction (via CNN), feature pooling (local and global feature sets using an attention mechanism), and prediction (pixel-level binary classification) [4]. By incorporating multiple frames, the model performs semantic segmentation and classifies lane types (e.g., dashed, solid), achieving state-of-the-art results on the VIL-100 dataset, demonstrating the value of temporal data for improved accuracy and classification [3, 5]. To optimize real-time performance, Wang and Li (2021) proposed a lightweight CNN based on UNet, with MobileNet as the encoder [6, 7, 8]. MobileNet's depthwise separable convolutions reduce parameters by separating channel-wise (depthwise) and pixel-wise (pointwise) operations, enabling efficient feature extraction [8]. The model achieves pixel-level segmentation at 40 FPS with fewer parameters, balancing performance and computational cost [6]. Similarly, de Sousa et al. (2021) developed a lightweight UNet-based model for embedded platforms, finding that reducing input image dimensions increased FPS by 35% but decreased accuracy by 27% [9].

Lane Detection Methods: Traditional methods in lane detection:-Traditional methods reviewed include geometric modeling (e.g., edge detection with Gaussian or Gabor filters, followed by Hough transform or B-spline curve fitting) and energy minimization techniques (e.g., conditional random fields, Kalman/particle filters for lane tracking) [15]. Deep learning methods encompass encoder-decoder CNNs (e.g., LaneNet, U-Net), fully convolutional networks with optimization (e.g., RANSAC), CNN+RNN approaches, and GAN-based models [16]. The proposed method distinguishes itself by treating lane detection as a time-series problem, using multiple frames for richer context, unlike single-frame approaches [7]. The project utilizes datasets like TuSimple for highways

and CULane for urban scenarios, supplemented by two new custom datasets capturing diverse driving conditions and rural roads [8]. It targets >95% accuracy, >0.3 Intersection over Union (IoU), and >60 frames per second (FPS) on edge devices [9]. By overcoming single-frame limitations and seamlessly integrating CNN and DRNN, this work advances real-time, scalable lane detection, contributing to safer and more reliable autonomous driving solutions [12]. Vision-Based Lane Detection Techniques:-Though the work targets at lane detection on unmarked roads, vision-based lane detection on marked, well-organized roads have been widely addressed in the literature [1], which provides guidelines to robust lane detection. These techniques rely on image processing and computer vision systems for detection of lane boundaries, predominantly used for purposes such as lane departure warning and lateral vehicle control in (ADAS). One such applications is by Mustafa Surti who uses the lanes Hough transform for lane detection [1]. It applies Canny edge detection in special Region of Interest (ROI) to keep track of the marked lane lines, presenting lane departure warning system on an embedded computer such as RaspberryPi [1]. Another approach involves a Gaussian function based road model and multi-stage image processing consisting of histogram computation, fitting and normalization, and polynomial function fitting for lane [2]. Curvature detection is also studied in another approach, in

this approach, images are transformed into the binary space and then divide into three partitions [3]. The curvature direction is indicated by the section with most non-zero pixels. One variation (Approach I) computes values of curve based on average histogram values of lower binary image, rather than the segmentation is used [3], [8]. [4] proposes an enhanced sliding window algorithm for lane curve fitting, which is robust to irregular lane markings. Curve value and vehicle lateral offset are calculated from road edge coordinates of binary images, and this method can improve the efficiency [4]. The BROGGI's GOLD system, which employs an edge-based lane border recognition algorithm, is another well-known system [9]. It maps the images into a bird-view, generates superpixels using adaptive filtering [9], and then concatenates the quasi-vertical brilliant lines into larger segments. The vehicle lateral position control output (u1) is obtained by comparing the template and the averaged scan line intensity profile taken from a portion of the upper row of the image, and by subsequently compensating the calculation for any lane curvature and offsets [10]. Further, the AURORA system, also developed by the same university, detects lane markings on structured roads with a camera in color pointing downwards attached to the side of a car [11]. These two methods have shown to be effective for well-marked roads, however, the difficulties seen when applying vision-based techniques for unmarked roads where no lane line existed, drive us to develop the new approach in the present work [12].

3. PROPOSED METHOD

Lane detection is one of the important task in ADAS which intends to detect solid or broken lines from the road surface for lane guidance in autonomous driving or for being assistance for driving human. Classic approaches such as geometric modeling (e.g., B-Snake [2] and Hough transform [3]) and semantic segmentation [16] generally analyze single images and do not work well in complex scenes, such as cities, low illumination and bad weathers [2]. These methods are not robust enough for practical ADAS, where consistent performance is required to ensure safety [12].

3.1 System Overview

The hybrid task in hand The proposed Hybrid network is designed to tackle such challenges by combining ability of DCNN to extract spatial information from single frames and, on the other hand capability of DRNN to model dynamical patterns over a few number of frames [7, 9]. These two mechanisms allow the system more robust for dynamic driving environments and enhance precision and robustness of the lane detection [8, 14].

3.2 Network Design

Architecture our proposed network model innovatively hybridize an encoder-decoder framework inspired by SegNet [8] and UNet [17] with a Convolutional Long Short- Term Memory (ConvLSTM) block [5]. The designed network is delicate between spatial data for each individual image and sequential data for video frames, thus it is suitable for advanced lane detection tasks. The architecture consists of two main parts: the LSTM network and the encoder-decoder network, which are designed separately for different purposes when detecting lanes.

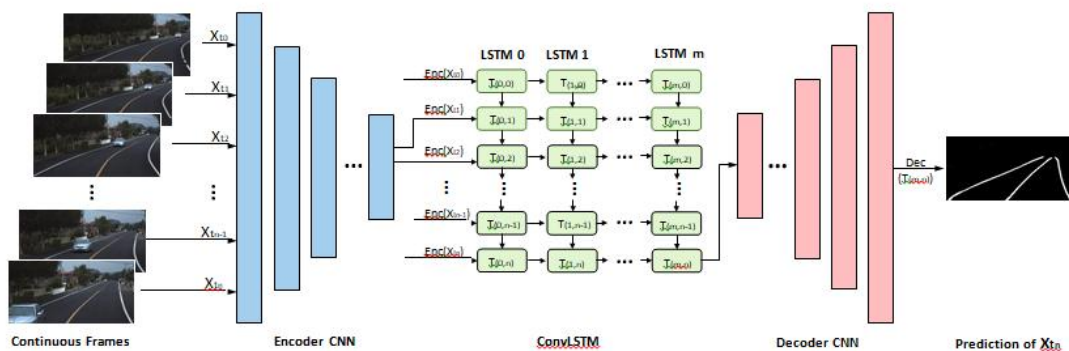


Figure 3.1 shows the proposed network's architecture.

3.3 Network LSTM:

The DRNN part takes recovered feature maps from the encoder CNN as inputs and treats consecutive frames of driving scenes in a time-series manner [9]. Such temporal modeling is important for estimating the continuity of lane patterns throughout the frames, in particular on dynamic scenarios [7].

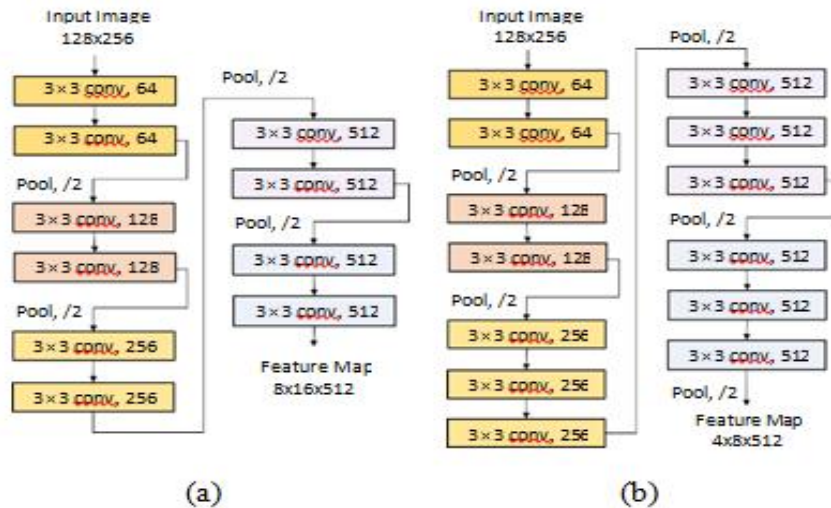
2) A double-layer Convolutional LSTM (ConvLSTM) is employed, which outperforms traditional Recurrent Neural Networks (RNNs) like standard LSTM or GRU [3]. Unlike traditional LSTM, which uses matrix multiplications, ConvLSTM applies convolution operations, reducing computational complexity and enabling efficient end-to-end training [15].

3) The ConvLSTM's strength lies in its ability to selectively "remember" important features (e.g., lane edges or colors) and "forget" irrelevant ones through its cell-based architecture, making it ideal for processing sequential video data [9]. The first layer extracts sequential features, while the second integrates them for cohesive temporal analysis [7].

4) By swapping out fully connected layers for convolutional operations, ConvLSTM significantly reduces both time and computational demands, enabling real-time performance for advanced driver assistance systems (ADAS) [5].

3.4 Encoder-Decoder Network:

The lane identification system described uses a semantic segmentation approach with an encoder-decoder architecture to detect lane markers at the pixel level, ensuring the output size matches the input for streamlined end-to-end training [16].



In Figure 3.2: the encoder networks (a) SegNet-ConvLSTM and (b) UNet-ConvLSTM are displayed.

The encoder processes images by extracting high-level features like lane shapes and textures through convolutional and pooling layers, inspired by SegNet's 16-layer VGGNet [6] and U-Net's design [7]. It reduces spatial dimensions while increasing feature depth and optimizes convolutional layers and kernels for accuracy and efficiency [7, 8, 13]. In the UNet-ConvLSTM and SegNet-ConvLSTM variants, the final encoder block limits kernel increases to simplify lane representation and support ConvLSTM processing [8]. The decoder reconstructs lane features using deconvolution and upsampling, with U-Net appending encoder-decoder feature maps to retain spatial details [7] and SegNet using stored encoding indices for precise recovery [8], ensuring accurate lane reconstruction [12, 14]. By integrating ConvLSTM into SegNet and U-Net, the system creates SegNet-ConvLSTM and UNet-ConvLSTM, combining spatial segmentation for static images with temporal learning for video sequences, enhancing lane detection for real-world driving [9, 7]. Convolutions employ a Convolution-BatchNorm-ReLU process with "same" padding to maintain feature map sizes [13, 16]. This architecture balances computational efficiency and detection accuracy, making it suitable for resource-constrained Advanced Driver Assistance Systems (ADAS) [6,8].

Optimization: The Adam optimizer [6] is used to update weights, leveraging its adaptive learning rate to ensure stable and efficient convergence during training [6].

End-to-End Training: The network's fully convolutional design, combined with ConvLSTM, supports seamless back-propagation across spatial and temporal components, enabling holistic optimization [16]. The training strategy leverages pre-trained weights to reduce computational demands while maintaining high accuracy, making the network practical for real-world deployment [4]. The use of ConvLSTM further enhances training efficiency by reducing the parameter complexity of traditional LSTM [3].

Segmentation Loss: Binary cross-entropy loss for classifying pixels as lane or non-lane.

Formula: $L_{seg} = -\sum [y \cdot \log(\hat{y}) + (1-y) \cdot \log(1-\hat{y})]$, where y is the ground truth and $\hat{y}$ is the predicted probability.

Discriminative Loss: Encourages clustering of features for the same lane and separation of different lanes.

Includes variance (within-lane compactness) and distance (between-lane separation) terms.

Polynomial Fitting Loss (optional): L2 loss to align predicted lane points with smooth polynomial curves.

The formula you provided, $L_{fit} = \sum(\hat{y} - p(x))^2$, is the polynomial fitting loss used in the lane detection model to ensure predicted lane points align with a smooth polynomial curve. Below, I clarify and paraphrase its definition and context within the CNN architecture for lane detection.

4. EXPERIMENTS AND RESULTS

4.1 Dataset:

The TuSimple dataset, a widely used benchmark for lane detection in autonomous driving, consists of 6,408 images organized into approximately 320 one-second video clips, each with 20 frames at a resolution of 1280x720 pixels [1]. Only the final (20th) frame in each clip is annotated with lane markings, which range from 2 to 5 lanes per image. The dataset captures diverse driving conditions, including varying weather (good to mediocre), daytime lighting, and traffic scenarios [1]. It is pre-split into 3,626 clips for training and validation and 2,782 clips for testing [1].

Lane annotations are provided in JSON format as one-pixel-thick polylines defined by X and Y coordinates [1]. For compatibility with deep learning segmentation models, these annotations are post-processed into binary segmentation masks—grayscale images matching the input image dimensions—where white pixels represent lane markers and black pixels indicate the background [2]. To improve training effectiveness, the thin polylines are thickened to 5 pixels in both dimensions, a choice based on trials balancing practical lane representation and fair evaluation [1]. The TuSimple evaluation protocol considers a predicted point correct if it lies within 20 pixels of an annotated point, supporting this thickening approach, though thicker polylines could enhance performance at the cost of unfair comparisons [1]. This preprocessing ensures the dataset is suitable for training and evaluating CNN-based segmentation models for real-time lane detection [2]. Evaluation metrics are critical due to the dataset's imbalance (few lane pixels compared to non-lane pixels), focusing on accuracy, robustness, and real-time computation [1, 2].

4.2 Parameter Analysis:-

1).Accuracy measures the percentage of correctly classified pixels (lane or non-lane) compared to ground truth [1]. While straightforward, it's less reliable for imbalanced datasets like those in lane detection, as it can be skewed by the dominance of non-lane pixels [2]. Despite this, accuracy is included for comparison with other studies, calculated as:

$$\text{Accuracy: } (TP + TN) / (\text{Total Number of Pixels}) \quad (1)$$

Precision is the ratio of true positives to (false positives) of the predicted lane pixels that are correct.

$$\text{Precision: } TP / (TP + FP) \quad (2)$$

Recall: Percentage of actual lane pixels correctly identified.

$$\text{Recall: } TP / (TP + FN) \quad (3)$$

F1 Score: Harmonic mean of precision and recall, providing a balanced measure of performance.

$$\text{F1 Score: } 2 * (\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (4)$$

4.3 Experiment:-

The benchmark experiment evaluates the efficiency and feasibility of a hybrid lane detection model that combines traditional vision techniques with deep learning, including a CNN and a DRNN with LSTM units, to enhance lane detection by leveraging spatial and temporal data [1].

The independent variables to which both the image unit model along with the temporal dynamics unit are appended, are the image-level features learned by the hybrid model and the temporal information across pairs of frames, which help in making the model become robust to eyeglass occlusions or shadowing [1]. The performance of the algorithm is evaluated using the following dependent variables: the detection performance of the algorithm is quantitatively analyzed using standard semantic segmentation metrics (accuracy, F1-Score) between 0 and 1, and inference time in Frames Per Second (FPS), an upper bounded positive integer to indicate the real-time potential [6].

The experiment tests whether the hybrid model achieves high accuracy (>95% on TuSimple), robust segmentation, and real-time performance (>60 FPS), providing a comprehensive comparison with current models to validate its suitability for autonomous driving [1, 3].

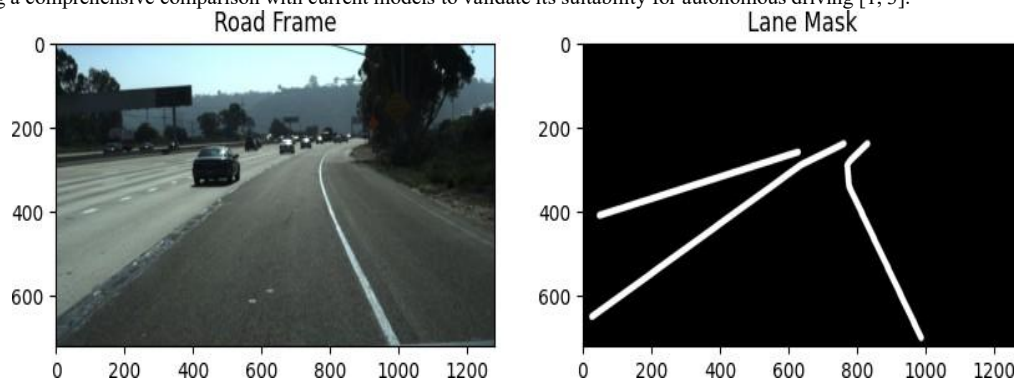


Figure 4.1: Labeled ground-truth lanes and an input image example. (a) The image input. (b) The ground truth.

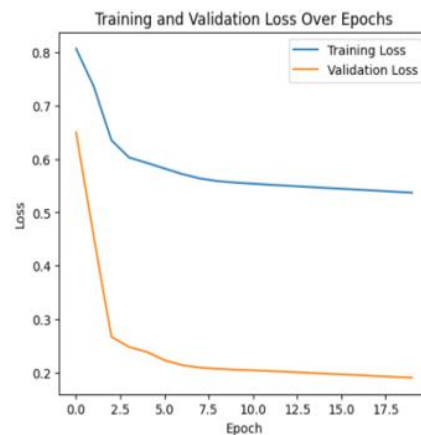


Figure4.2: Training and Validation Loss

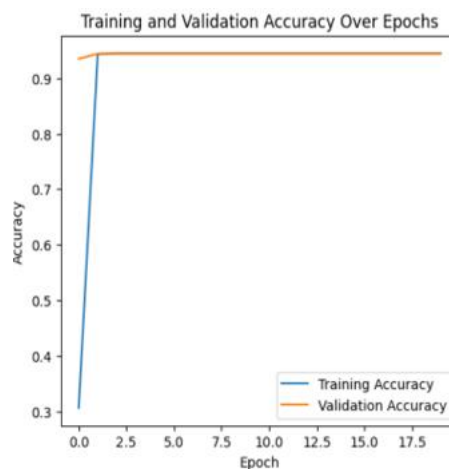


Figure4.3: Training and Validation Accuracy

CONCLUSION

The paper presents a novel hybrid neural network for robust lane detection, integrating a CNN and a RNN with a Convolutional Long Short-Term Memory (ConvLSTM) unit within an encoder-decoder framework. This architecture processes multiple continuous frames to predict lane markings in the current frame via semantic segmentation, offering improved performance over single-frame methods. The CNN encoder extracts features from each frame, the ConvLSTM processes these as a temporal sequence, and the CNN decoder reconstructs the information to produce lane predictions. Two custom datasets of continuous driving scenes were created to evaluate the model's performance.

Compared to baseline models using single images, the proposed architecture achieves significantly better results, demonstrating the advantage of leveraging temporal information from multiple frames. The use of ConvLSTM outperforms traditional Fully Connected LSTM (FcLSTM) in sequential feature learning and lane prediction, yielding higher precision, recall, and accuracy. Testing on a challenging dataset with diverse driving conditions (e.g., occlusions, shadows) confirmed the model's robustness and ability to minimize false positives. Analysis showed that longer input sequences further enhance performance, reinforcing the value of multi-frame inputs.

Future work includes integrating lane fitting into the framework for smoother, more cohesive lane boundaries and investigating why SegNet-ConvLSTM outperforms UNet-ConvLSTM in dim environments with strong interferences, to further optimize the model. The proposed hybrid model advances real-time, reliable lane detection for autonomous driving applications.

REFERENCES

- [1] W. Wang, D. Zhao, W. Han, and J. Xi, "A learning-based approach for lane departure warning systems with a personalized driver model," *IEEE Transactions on Vehicular Technology*, 2018.
- [2] Y. Wang, K. Teoh, and D. Shen, "Lane detection and tracking using b-snake," *Image and Vision computing*, vol. 22, pp. 269–280, 2004.
- [3] A. Borkar, M. Hayes, and M. Smith, "Polar randomized hough transform for lane detection using loose constraints of parallel lines," in *International Conference on Acoustics, Speech and Signal Processing*, 2011.
- [4] C. Wojek and B. Schiele, "A dynamic conditional random field model for joint labeling of object and scene classes," in *European Conference on Computer Vision (ECCV)*, 2008.
- [5] J. Hur, S.-N. Kang, and S.-W. Seo, "Multi-lane detection in urban driving environments using conditional random fields," in *IEEE Intelligent Vehicles Symposium (IV)*, 2013, pp. 1297–1302.
- [6] J. Kim, J. Kim, G.-J. Jang, and M. Lee, "Fast learning method for convolutional neural networks using extreme learning machine and its application to lane detection," *Neural networks*, vol. 87, pp. 109–121, 2017.

- [7] J. Li, X. Mei, D. V. Prokhorov, and D. Tao, "Deep neural network for structural prediction and lane detection in traffic scene," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 28, pp. 690–703, 2017.
- [8] S. Lee, J.-S. Kim, J. S. Yoon, S. Shin, O. Bailo, N. Kim, T.-H. Lee, H. S. Hong, S.-H. Han, and I.-S. Kweon, "VPGNet: Vanishing point guided network for lane and road marking detection and recognition," in *International Conference on Computer Vision*, 2017, pp. 1965–1973.
- [9] Y. Huang, S. tao Chen, Y. Chen, Z. Jian, and N. Zheng, "Spatial-temporal based lane detection using deep learning," in *IFIP International Conference on Artificial Intelligence Applications and Innovations*, 2018.
- [10] R. B. Girshick, "Fast r-cnn," in *IEEE International Conference on Computer Vision (ICCV)*, 2015, pp. 1440–1448.
- [11] J. Redmon, S. K. Divvala, R. B. Girshick, and A. Farhadi, "You only look once: Unified, real-time object detection," in *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 779–788.
- [12] K. He, G. Gkioxari, P. Dollár, and R. B. Girshick, "Mask R-CNN," in *International Conference on Computer Vision*, 2017, pp. 2980–2988.
- [13] J. LTE-A heterogeneous networks using femtocells, *International Journal of Innovative Technology and Exploring Engineering*, 2019, 8(4), pp. 131–134 (SCOPUS) Scopus cite Score 0.6
- [14] Power Control Schemes for Interference Management in LTE-Advanced Heterogeneous Networks, *International Journal of Recent Technology and Engineering (IJRTE)* ISSN: 2277-3878, Volume-8 Issue-4, November 2019, pp. 378-383 (SCOPUS)
- [15] Performance Analysis of Resource Scheduling Techniques in Homogeneous and Heterogeneous Small Cell LTE-A Networks, *Wireless Personal Communications*, 2020, 112(4), pp. 2393–2422 (SCIE) {Five year impact factor 1.8 (2022)} 2022 IF 2.2, Scopus cite Score 4.5
- [16] Victim Aware AP-PF CoMP Clustering for Resource Allocation in Ultra-Dense Heterogeneous Small-Cell Networks. *Wireless Personal Commun.* 116(3): pp. 2435-2464 (2021) (SCIE) {Five-year impact factor 1.8 (2022)} 2022 IF 2.2, Scopus cite Score 4.5