

Vehicle recognition on indian roads using data augmentation and VGG-16 model

Arunkumar K. L.¹, Poornima K. M.², Ajit Danti³, Manjunatha H. T.¹

¹Department of Computer Applications, J N N College of Engineering (JNNCE), Visvesvaraya Technological University (VTU), Belagavi, India

²Department of Computer Science and Engineering, J N N College of Engineering (JNNCE), Visvesvaraya Technological University (VTU), Belagavi, India

³Department of Computer Science and Engineering, School of Engineering and Technology Christ (Deemed to be university), Kengeri Campus, Bangalore, India

Article Info

Article history:

Received Feb 14, 2024

Revised Aug 29, 2025

Accepted Oct 15, 2025

Keywords:

CMMR

CNN

SIFT

SVM

VGG16

VMMR

ABSTRACT

In an advanced intelligent transportation system vehicle recognition and classification is very significant. In current research trend, recognition of vehicles is done by using machine learning (ML) and computer vision techniques. Vehicle's multi-view images or videos with different lighting conditions are annotated and given to the deep neural network to build an automated system to recognize the vehicles models. The augmentation of data can increase the number of samples in learning, with the small available datasets. Geometric transformations, brightness changes, and different filter operations are applied to the data through data augmentation. Furthermore, be orthogonal experiments we determine the optimal data augmentation method to obtain 96% accuracy in results. Detailed information is reported based on the classification of four different types of vehicles and the results show that convolutional neural network with 16 layers deep techniques are effective in solving challenging tasks while recognizing moving vehicles.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Arunkumar K. L.

Department of Computer Applications, JNN College of Engineering (JNNCE)

Visvesvaraya Technological University (VTU)

Belagavi, Karnataka, India

Email: alihhas@upm.edu.my

1. INTRODUCTION

In the era of smart transportation and artificial intelligence (AI), the application of computer vision techniques to recognize Indian car makes and models represents a groundbreaking advancement with significant implications for various industries. This introduction provides an overview of the importance, challenges, and potential solutions in implementing a computer vision system tailored to identify and classify Indian car makes and models. The Indian automotive landscape is characterized by a diverse range of car makes and models, each catering to specific market segments and consumer preferences. Efficiently recognizing and classifying these vehicles can offer a myriad of benefits, including improved traffic management, enhanced security, and valuable insights for businesses such as insurance and retail.

The complexities of the Indian automotive environment poses challenges for accurate car make and model recognition. Factors such as variations in vehicle design, diverse weather conditions, and the presence of customized or locally modified vehicles necessitate sophisticated computer vision algorithms capable of

handling this unique diversity. Computer vision, a field within AI, enables machines to interpret and make decisions based on visual data. By leveraging techniques such as image processing, feature extraction, and deep learning, computer vision systems can be trained to recognize intricate details that distinguish one car make and model from another. AI and Machine learning (ML) is having tremendous scope in this digital era. Hence, researchers will have an opportunity to work with advanced techniques. In recognition of objects using ML, convolution neural network (CNN) is widely used in most of the recent research work. In deep learning the amount of pre-processing task will be less while compared with other classification techniques.

For recognition of objects in natural scenes, CNNs are one of the best types of neural networks for image processing [1]. CNN is predominantly used in many research works on face recognition and traffic sign board detections. In CNN standard algorithms can be used to perform the task by applying suitable filters and spatial and temporal dependencies with low cost and without using large resources [2]. CNNs, which are inspired by visual cortex, are similar to neurons in the human brain.

2. LITERATURE SURVEY

Tremendous of research is going on vehicle make and model recognition system. I referred research papers on VMMR employing deep learning and ML in image processing for my work. Jain and Kumar [3] uses the logo which appears in front and rear view of the vehicle to recognize the maker of the vehicle. They scale invariant feature transform (SIFT) key points to localize the vehicle logo. Later features are extracted and feature matrix is created for training samples. Later template matching is performed for the logo features extracted in testing sample. They have reported around 70

Petrovic and Cootes [4] used licence plate as a region of interest to extract the features and matching. They used nearest neighbour algorithm for classification of characters in licence plate. Similarly, Cheung and Chue [5] in the paper “Analysis of features for rigid structure vehicle type recognition” used two methods for feature extraction and matching. They also used SIFT key points as features and descriptors for matching the query car image with its database images. The second approach used was Harris Corners for interest point detection and fast normalized correlation for feature matching.

In the bag-of-features (BoF) approach used by Siddiqui *et al.* in their paper “Real-Time Vehicle Make and Model Recognition based on a Bag-of-SURF-Features” [6], used the BoF model using SURF feature extractor. A vocabulary of words or features was formed by clustering SURF features of the images, next the classification was done using support vector machines (SVM) classifier.

Zhou *et al.* [7], in their paper “Image based Vehicle Analysis using Deep Neural Nets...” used a pre-trained neural network ‘AlexNet’. Although, the AlexNet originally is trained to recognize 1,000 categories of different objects, authors used it to distinguish between various categories of cars. Two approaches were used, the first one was to extract features from higher layers of the CNN and then classify using SVM classifiers. The second approach was to fine tune the neural network to be used for VMMR system. Another paper named “View Independent Car Make and Model and Color Recognition” by Dehgan *et al.* [8], authors trained a deep neural network for VMMR using a very large image data set consisting of millions of images. Their system is known as “Sighthound”. Different approaches are made to recognise the vehicles and different traffic elements on Indian roads using computer vision techniques in my early work [9]-[22].

3. METHOD

Proposed system architecture is shown in Figure 1. It consists of different stages. Experimental factors are determined in first step along with choosing the proper levels. The professional knowledge and experience help in choosing the experimental factors [23]. Factors include translation, brightness, rotation, image zoom, and four different filter procedures used as common data augmentation methods. For each factor only three kinds of levels are chosen at the experimental level, at the run time many levels are increased.

3.1. Image acquisition

Three different databases of images were used to create the image data sets we used. The first was the CompCars database which is a comprehensive cars picture database consists of images from surveillance cameras, and the second database having images from internet resources. The third data set was constructed by us. Approximately 10,000 images were collected in total and images were first divided into two groups training and testing images.



Figure 1. The workflow of our proposed VMMR system architecture

3.2. Data preprocessing

Segmented images are employed for the purpose of data augmentation. The data augmentation technique is utilized on the training images to augment the number of image samples in the dataset. By employing the data augmentation method, the quantity of training samples in the dataset can be expanded. To reduce the time required for experiments, the factors are divided into two parts: one pertains to the brightness of the image, while the other encompasses geometric transformations and image filters. Various geometric transformations, such as rotation, translation, scaling, and brightness adjustments, are applied to the segmented image. Additionally, as part of the data augmentation experiments, different filter operations like bilateral filters are applied to the same segmented image. These augmented datasets are subsequently utilized to establish the training dataset. Table 1 presents each data augmentation level along with the methods and parameters employed to achieve the data augmentation.

Table 1. Experiments with image brightness and orthogonal geometric transformation

Methods	Parameters	Level 1	Image Level 2	Level 3
Image zoom	Scale coefficient	1.2	1.5	1.8
Image translation	$(\Delta x, \Delta y)$ in Equation (1)	(12,10)	(18,15)	(24,20)
Image rotation	Rotation angle θ	10	25	40
Image brightness	Scale coefficient	1.2	1.5	1.8
Mean filter	Window size	3×3	4×4	5×5
Median filter	Window size	3×3	5×5	7×7
Gaussian filter	Window size	3×3	5×5	7×7
Bilateral filter	Window size	3×3	5×5	7×7

3.3. Establishing training set

In a deep learning, data labelling is a necessary step. To establish the training dataset before feeding into the CNN model annotation is very important. Annotation of images can be done in different ways like bounding box, polynomial segmentation, 3D cuboids, key-point and Landmark, Lines and Splines many more. In computer vision the most common type of annotation method is drawing a bounding box for the object to be trained in CNN model. A rectangular box is drawn for the target object. This helps to find the location of the target object using x-axis and y-axis coordinates of the top left corner and bottom right corner of the bounding box. These bounding boxes are generally used for localization and object detection. Like this, annotations are done on the augmented images and these set of images are used to train the system for recognition of vehicle make and model.

3.4. CNN model

Car make and model recognition (CMMR) can accomplished by modifying the visual geometry group (VGG-16) network for pretrained models. Fine-tuning is a technique for adjusting the parameters of a CNN model for a specific region in the pretrained dataset [24]. In the Figure 2, the input RGB image size is adjusted to 224×224 for the network. This deep CNN's architecture is consisting of 13 convolutional layers, 5 pooling layers, and 3 fully connected layers, each with 54 channels. The depth of network is increased in VGG-16 architecture by appending many convolution layers with kernel size 3×3 and these kernels with small size perform well in successfully completion of work. With the size of convolution stride of 1 and spatial padding of 1, we can preserve the spatial resolution. Five max pool layers are used for down sampling, pooling layers are followed by some convolutional layers. With stride of 2 and with pixel window of size 2×2 , max pooling layers are achieved. Convolution layers and fully connected-layers utilises rectified linear unit (ReLU) activation functions and in the last layer utilised the SoftMax function.



Figure 2. An illustration of CNN model architecture

In the train model, a training dataset is used, and a forward propagation neural network with multiple layers is used to compute the output. The value returned from the ReLU function is assigned to C and using this features are mapped in the convolution layer.

$$C = \varphi(H(x, y)) \quad (1)$$

Where the ReLU function is represented by $\varphi(\cdot)$.

$$\varphi(H(x, y)) = \max(0, H(x, y)) \quad (2)$$

$$H(x, y) = \sum_{m, n \in s}^l W(m, n) l'_i(x + m, y + n) + b$$

Where weight of the kernel matrix is denoted by W and b the bias $\varphi(\cdot)$ represents activation function and H indicates sigmoid; Results of the ReLU function accelerates network convergence and reduces computation when compared with different activation functions like tanh and sigmoid. Along with this, gradient of sigmoid or tanh functions approaches zero as the absolute value of $H(x, y)$ increases. Hence activation of these functions will lead to vanish the gradient problem. But this is not a problematic in ReLU when $H(x, y) > 0$.

Next to the convolution layer we have max pooling layer with kernel of 2×2 . In order to reduce the number of parameters and network spatial size, this pooling layer is utilised. This in turn helps to avoid overfitting. P denotes the pooling layer's output map. Calculation of P is as follows.

$$P = g(C) \quad (3)$$

Where $g()$ is used to find the max value. This function selects the maximum value in the window and rejects the other values, while windows move across C . To ease the overfitting dropout layers are utilized. During each training phase, 0 is set to neuron at the output with a probability of 0.5. These neurons will not contribute for forward propagation and back-propagation, this helps in prevention of overfitting and also these neurons are called "dropout".

$$F_q = \varphi\left(\sum_{m, n \in s}^l W(m, n) P(x, y) + b\right) \quad (4)$$

In the network's loss function L is used to denote the SoftMax-loss function and using mini-batch gradient method our model is trained.

$$L = -\frac{1}{m} \varphi \left(\sum_{i=1}^M \sum_q^j T_i \log(P_q) \right) \quad (5)$$

Where a batch of one iteration, number of images are represented by M and network output at neuron q is represented by p_q , by using the SoftMax function we can calculate the probability of the model's prediction,

$$P_q = \frac{\exp(F_q)}{\sum_{z=1}^j \exp(F_z)} \quad (6)$$

weight is denoted by w and is updated using back-propagation algorithm, Z indicates number of batches and J indicated the number of classes. Rule to update w is as follows.

$$\Delta v_t = \mu v_{t-1} - \alpha \frac{\delta L}{\delta \omega'} \quad (7)$$

$$w'_t = w_t + \Delta v_t \quad (8)$$

To accelerate the convergence, momentum coefficient is used which is denoted by μ ($\mu = 0.9$), value of the weight which is previously updated is denoted by Δv_{t-1} , at every iteration t current weight is denoted by w_t , and learning rate is denoted by α , 0.001 is the basic learning rate of α_0 is 0.001. The rule to update α is,

$$\alpha_j = \alpha_{j-1} * \gamma^{\frac{t}{u}} \quad (9)$$

where gamma parameters are denoted by γ ($\gamma = 0.1$) and stepsize denoted by u (u is 10,000) algorithm 1 explains working nature of our proposed method. Training sample pictures are acquired using VMMR algorithm. To increase the number of training samples data augmentation is used such as geometric transformation, image filter operations or brightness manipulations. Using these augmented training samples the CNN model is trained. At the time of testing, untrained samples given to the trained model to recognize the vehicle make and model.

Algorithm 1. Car make and model recognition

1. Samples of vehicle pictures are captured;
2. The training dataset is established by augmenting the samples in training folder;
3. **loop** foreach iteration from t to $t_{\max}$, **do**
4. if $t\%$ value of stepsize is 0 **then** (t is defined in equation (7))
5. α is updated which indicates the rate of learning; (α is defined in equation (8))
6. **condition termination**
7. loop for training every image I_0 to k **do**
8. Using forward propagation algorithm output of each layer is computed;
9. **end for**
10. Using back propagation algorithm weight ω is updated
11. **end for**
12. **for** each image in test folder **do**
13. Using forward propagation algorithm output of each layer is computed;
14. Computed the prediction result as ρ ;
15. Display the Vehicle make and model name;
16. **end for**

3.5. Testing

The performance of the VMMR is demonstrated and evaluated using 5-fold cross-validation. The effect of different augmentation methods is analysed using range analysis method. Range analysis method is more spontaneous compared with other analysis methods. R is used to denote the range value, which indicates the importance of the corresponding factor. Finally, to predict the vehicle make and model a custom model is generated by training with the CNN.

4. EXPERIMENTAL RESULTS

For experimentation purpose around 3,750 images are trained, which contributed to make a total of 4 classes, or categories of cars. In this supervised learning method sample images are annotated before training. Experimental output is shown in Figure 3. Assuming that each class comprises J vehicles, the collection of original car photos is defined by $I = I_k | k[1, N]$, where I_k signifies input car image and total number of original images are denoted by N . To augment the original training samples best data augmentation techniques are used. The performance of our method is demonstrated and using 5-fold cross-validation is used for performance evaluation. Experimental output is shown in Figure 3. Figure 3(a) shows image classification and Figure 3(b) shows vehicle recognition.

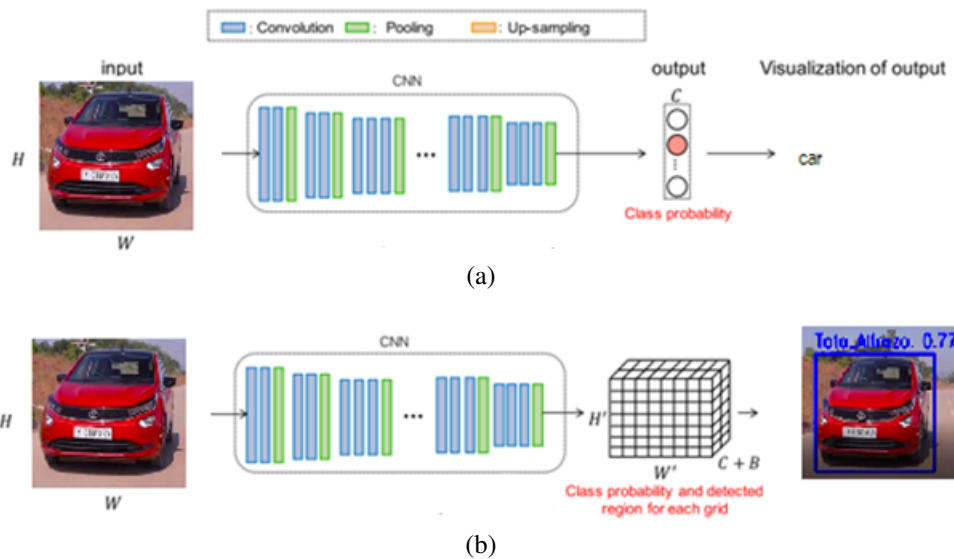


Figure 3. Experimental output (a) image classification and (b) vehicle recognition

Four folds are used for model training and remaining folds are used for testing. 83.5% is the average accuracy obtained for 5-fold cross-validation is shown in Figure 4. These findings demonstrate that augmented data combined with a deep CNN model can significantly increase VMMR accuracy when training samples are minimal are shown in Table 2 and Table 3.

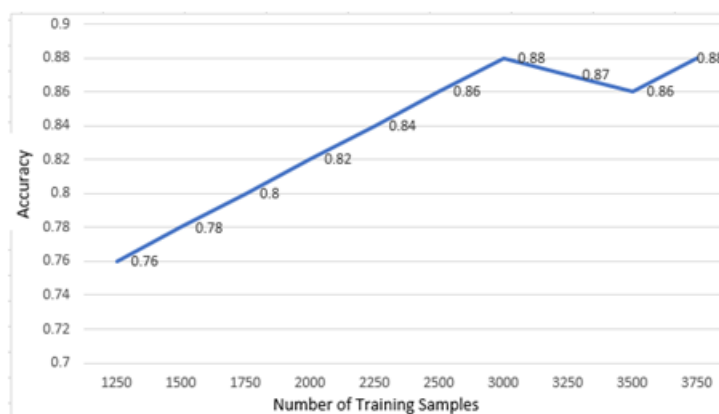


Figure 4. The recognition performance of different number of training samples

In Table 2, we present a comparison of various combinations of levels and factors related to data augmentation methods. Here, K1 represents the sum of the l th current factor, while k_1 signifies the average of K1. The data in Table 2 highlights that the highest range value corresponds to image rotation, indicating its significance as the primary factor. Furthermore, the impact of factors is arranged in descending order within the table. Based on the information provided, it is evident that the optimal combination for geometric transformation and image brightness involves utilizing the 3rd level of image zoom, the 1st level of image translation, the 1st level of image rotation, and the 3rd level of image brightness.

Table 2. Experiments with image brightness and orthogonal geometric transformation

Levels	Image zoom	Image translation	Image rotation	Image brightness	Accuracy
1	1	1	1	1	83.3%
2	1	2	2	2	79.6%
3	1	3	3	3	81.5%
4	2	1	2	3	83.3%
5	2	2	3	1	83.3%
6	2	3	1	2	83.3%
7	3	1	3	2	83.3%
8	3	2	1	3	85.2%
9	3	3	2	1	81.5%
K1	244.4	249.9	251.8	248.1	
K2	249.9	248.1	244.4	246.2	
K3	250	246.3	248.1	250.0	
$\bar{k}_1 = \frac{K_1}{3}$	81.47	83.30	83.93	82.70	
$\bar{k}_2 = \frac{K_2}{3}$	83.30	82.70	81.47	82.07	
$\bar{k}_3 = \frac{K_3}{3}$	83.33	82.10	82.70	83.33	
R	1.86	1.20	2.46	1.26	

The outcomes of orthogonal experiments conducted on filter operation are detailed in Table 3. It is evident that the bilateral filter exhibits the widest range, indicating its significant impact. To assess the impact of bilateral filter and image translation factors, we applied both to enhance the same original samples and compared the accuracy of vehicle recognition. The accuracy achieved when training the samples with bilateral filter is 74.1%, while the accuracy with image translation is 79.6%. When it comes to data augmentation, it is advisable to opt for the factor that yields better results. Consequently, the most effective data augmentation method comprises the 3rd level of image zoom, 1st level of image translation, 1st level of image rotation, and 3rd level of image brightness.

Table 3. Filter operation using orthogonal experiments

Levels	Mean filter	Median filter	Gaussian filter	Bilateral filter	Accuracy
1	1	1	1	1	81.5%
2	1	2	2	2	77.8%
3	1	3	3	3	79.6%
4	2	1	2	3	79.6%
5	2	2	3	1	85.2%
6	2	3	1	2	77.8%
7	3	1	3	2	79.6%
8	3	2	1	3	81.5%
9	3	3	2	1	79.6%
K1	238.9	240.7	240.8	246.3	
K2	242.6	244.5	237	235.2	
K3	240.7	237	244.4	240.7	
$\bar{k}_1 = \frac{K_1}{3}$	79.63	80.23	80.27	82.10	
$\bar{k}_2 = \frac{K_2}{3}$	80.87	81.50	79.00	78.40	
$\bar{k}_3 = \frac{K_3}{3}$	80.23	79.00	81.47	80.23	
R	1.24	2.50	2.47	3.70	

After, the best data augmentation method is used to augment the original training samples. To demonstrate the performance of our method, the performance of our method is evaluated using 5-fold cross-validation. The model is trained on four folds, and tested on the remaining fold. The average accuracy of 5-fold cross-validation is 86.3%. The result shows that using the deep CNN model with data augmentation can effectively improve the accuracy of vehicles recognition based on a small number of training samples.

5. DISCUSSION

After, the best data augmentation method is used to augment the original training samples. To demonstrate the performance of our method, the performance of our method is evaluated using 5-fold cross-validation. The model is trained on four folds, and tested on the remaining fold. The average accuracy of 5-fold cross-validation is 86.3%. The result shows that using the deep CNN model with data augmentation can effectively improve the accuracy of vehicles recognition based on a small number of training samples. As shown in Figure 4, more samples are used for training; the model gives the higher the accuracy and better performance. Additionally, with 3,750 samples, we can achieve an accuracy of 98.1

6. CONCLUSION

A novel method is proposed for CMMR with augmented data using CNN model. Various augmentation methods performance is analysed based on the orthogonal experiments. The best augmentation method which is suitable for the proposed work is determined using orthogonal table. Later, we demonstrated the accuracy 83.5% obtained in recognition of vehicle using augmented data with deep learning. If input videos captured in the better lighting conditions, the accuracy of our method could improve to higher rates.

FUNDING INFORMATION

Authors state no funding involved.

CONFLICT OF INTEREST STATEMENT

Authors state no conflict of interest.

DATA AVAILABILITY

Data availability is not applicable to this paper as no new data were created or analyzed in this study.

REFERENCES

- [1] S. K. Roy, G. Krishna, S. R. Dubey, and B. B. Chaudhuri, "HybridSN: exploring 3-D-2-D CNN feature hierarchy for hyperspectral image classification," *IEEE Geoscience and Remote Sensing Letters*, vol. 17, no. 2, pp. 277–281, Feb. 2020, doi: 10.1109/LGRS.2019.2918719.
- [2] R. He, X. Wu, Z. Sun, and T. Tan, "Wasserstein CNN: learning invariant features for NIR-VIS face recognition," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 41, no. 7, pp. 1761–1773, Jul. 2019, doi: 10.1109/TPAMI.2018.2842770.
- [3] M. Jain and D. T. Kumar, "Car make and model recognition." IIT Hyderabad, ODF, Yeddumailaram – 502205, 2015.
- [4] V. S. Petrovic and T. Cootes, "Analysis of features for rigid structure vehicle type recognition," in *Proceedings of the British Machine Vision Conference 2004*, 2004, pp. 61.1-61.10, doi: 10.5244/C.18.61.
- [5] S. Cheung and A. Chu, "Make and model recognition of cars," 2008, [Online]. Available: <http://cseweb.ucsd.edu/classes/wi08/cse190-a/reports/scheung.pdf>.
- [6] A. J. Siddiqui, A. Mammeri, and A. Boukerche, "Real-time vehicle make and model recognition based on a bag of SURF features," *IEEE Transactions on Intelligent Transportation Systems*, vol. 17, no. 11, pp. 3205–3219, Nov. 2016, doi: 10.1109/TITS.2016.2545640.
- [7] Y. Zhou, H. Nejati, T.-T. Do, N.-M. Cheung, and L. Cheah, "Image-based vehicle analysis using deep neural network: a systematic study," in *2016 IEEE International Conference on Digital Signal Processing (DSP)*, Oct. 2016, pp. 276–280, doi: 10.1109/ICDSP.2016.7868561.
- [8] A. Dehghan, S. Z. Masood, G. Shu, and E. G. Ortiz, "View independent vehicle make, model and color recognition using convolutional neural network," *arXiv preprint arXiv:1702.01721*, 2017, [Online]. Available: <http://arxiv.org/abs/1702.01721>.
- [9] C. Wang and Z. Peng, "Design and implementation of an object detection system using faster R-CNN," in *2019 International Conference on Robots & Intelligent System (ICRIS)*, Jun. 2019, pp. 204–206, doi: 10.1109/ICRIS.2019.00060.

- [10] L. Qian, Y. Fu, and T. Liu, "An efficient model compression method for CNN based object detection," in *2018 IEEE 9th International Conference on Software Engineering and Service Science (ICSESS)*, Nov. 2018, vol. 2018-Novem, pp. 766–769, doi: 10.1109/ICSESS.2018.8663809.
- [11] W. Zhang, J. Li, and S. Qi, "Object detection in aerial images based on cascaded CNN," in *2018 International Conference on Sensor Networks and Signal Processing (SNSP)*, Oct. 2018, pp. 434–439, doi: 10.1109/SNSP.2018.00088.
- [12] P. Pooyoi, P. Borwonginn, J. H. Haga, and W. Kusakunniran, "Snow scene segmentation using CNN-based approach with transfer learning," in *2019 16th International Conference on Electrical Engineering/Electronics, Computer, Telecommunications and Information Technology (ECTI-CON)*, Jul. 2019, pp. 97–100, doi: 10.1109/ECTI-CON47248.2019.8955140.
- [13] K. L. Arunkumar, A. Danti, and H. T. Manjunatha, "Classification of vehicle make based on geometric features and appearance-based attributes under complex background," in *Communications in Computer and Information Science*, vol. 1035, 2019, pp. 41–48.
- [14] K. L. Arunkumar and A. Danti, "A novel approach for vehicle recognition based on the tail lights geometrical features in the night vision," *International Journal of Computer Engineering and Applications*, 2018.
- [15] H. T. Manjunatha, A. Danti, and K. L. Arunkumar, "A novel approach for detection and recognition of traffic signs for automatic driver assistance system under cluttered background," in *Communications in Computer and Information Science*, vol. 1035, 2019, pp. 407–419.
- [16] K. L. Arunkumar, A. Danti, and H. T. Manjunatha, "Estimation of vehicle distance based on feature points using monocular vision," in *2019 International Conference on Data Science and Communication (IconDSC)*, Mar. 2019, pp. 1–5, doi: 10.1109/IconDSC.2019.8816996.
- [17] K. L. Arunkumar, A. Danti, H. T. Manjunatha, and D. Rohith, "Classification of vehicle type on indian road scene based on deep learning," in *Communications in Computer and Information Science*, vol. 1380, 2021, pp. 183–192.
- [18] H. T. Manjunatha, A. Danti, K. L. A. Kumar, and D. Rohith, "Indian road lanes detection based on regression and clustering using video processing techniques," in *Communications in Computer and Information Science*, vol. 1380, 2021, pp. 193–206.
- [19] K. L. Arunkumar and A. Danti, "Recognition of vehicle using geometrical features of a tail light in the night vision," *National conference on computation science and soft computing (NCCSSC-2018)*, 2018.
- [20] H. T. Manjunatha and A. Danti, "Detection and classification of potholes in indian roads using wavelet based energy modules," in *2019 International Conference on Data Science and Communication (IconDSC)*, Mar. 2019, pp. 1–7, doi: 10.1109/IconDSC.2019.8816878.
- [21] H. T. Manjunatha and A. Danti, "Indian traffic sign board recognition using normalized correlation method," *International Journal of Computer Engineering and Applications*, vol. XII, No. III, ISSN: 2321-3469.
- [22] H. T. Manjunatha and A. Danti, "Segmentation of traffic sign board in a cluttered background using using digital image processing," *National Conference on Network Security, Image Processing and Information Technology*, 2017.
- [23] Z. Pei, Y. Zhang, Z. Lin, H. Zhou, and H. Wang, "A method of image processing algorithm evaluation based on orthogonal experimental design," in *2009 Fifth International Conference on Image and Graphics*, Sep. 2009, pp. 629–633, doi: 10.1109/ICIG.2009.115.
- [24] W. Ouyang, X. Wang, C. Zhang, and X. Yang, "Factors in finetuning deep model for object detection with long-tail distribution," in *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, Jun. 2016, vol. 2016-Decem, pp. 864–873, doi: 10.1109/CVPR.2016.100.

BIOGRAPHIES OF AUTHORS



Arunkumar K. L.     is assistant professor at Department of computer applications, JNN College of Engineering, Visvesvara Technological University, Belagavi, Karnataka, India. He holds a masters degree in Computer Applications with specialization in image analysis. His research areas are image processing, biometrics, medical image analysis and pattern recognition. He is a recipient of different national and international awards such as NUST overall best researcher award. His research interests include image/signal processing, biometrics, medical image and analysis, and pattern recognition. He can be contacted at email: arunkumarkl@jnnce.ac.in.



Poornima K. M.     received the B.Sc. degree in computer science from the University of Hartford, USA, the M.Sc. degree in computing from the University of Bradford, U.K., and the Ph.D. degree in soft computing from The University of Sheffield, U.K. She used to hold several administrative posts with the School of Computing, Universiti Teknologi Malaysia (UTM), Johor, from 2007 to 2018, including the Head of Department, the Deputy Dean of Postgraduate Studies, and the Deputy Dean of Academic. She was also the Director of the Big Data Centre (Centre of Excellence), UTM, from 2019 to February 2020. She is currently a professor with the Department of Data Science and the Dean of Undergraduates, Universiti Malaysia Kelantan (UMK). She has supervised and co-supervised more than 20 masters and 20 Ph.D. students. She has authored or coauthored more than 150 publications: 80 proceedings and 57 journals, with 19 H-index and more than 1000 citations. Her research interests include soft computing, ML, and intelligent systems. She can be contacted at email: kmpoornima@jnnce.ac.in.



Ajit Danti     his higher studies in MS Food Engineering by coursework at the at the department of Food Science, University of Leeds from 1981-1982. He is attached to the Faculty of Biotechnology and Biomolecular Sciences. His research area then was on spray drying of food. With a small research grant provided by UPM, he developed the process for producing spry-dried coconut milk which made the national headlines. His vast experience and expertise in the field of biotechnology and biomolecular sciences have enabled him to become a national point of reference in the area of biomass, renewable energy and waste utilization. He has also served as a consultant to The Science Advisor Office, Prime Minister's Department, on the national project on biomass utilisation, and is the national representative for the Asia Biomass Association headquartered in Tokyo, Japan. He can be contacted at email: ajit.danti@christuniversity.in.



Manjunatha H. T.     his higher studies in MS Food Engineering by coursework at the at the department of Food Science, University of Leeds from 1981-1982. He is attached to the Faculty of Biotechnology and Biomolecular Sciences. His research area then was on spray drying of food. With a small research grant provided by UPM, he developed the process for producing spry-dried coconut milk which made the national headlines. His vast experience and expertise in the field of biotechnology and biomolecular sciences have enabled him to become a national point of reference in the area of biomass, renewable energy and waste utilization. He has also served as a consultant to The Science Advisor Office, Prime Minister's Department, on the national project on biomass utilisation, and is the national representative for the Asia Biomass Association headquartered in Tokyo, Japan. He can be contacted at email: manjudeepa@jnnce.ac.in.