

DEEP LEARNING MODEL FOR TRACKING HUMAN BEHAVIOR IN SUPERMARKET THROUGH CCTV FOOTAGE

Aarthi Gopalakrishnan^{1,*}, A. Neelima² and C. Priya³

¹Research Scholar, Department of Computer Science and Engineering, B.S. Abdur Rahman Crescent Institute of Science and Technology, Chennai, India.

²Department of Computer Science and Engineering, Koneru Lakshmaiah Education Foundation, Vaddeswaram, Andhra Pradesh, India.

³Lecturer, Department of Computer Engineering, Government Polytechnic College, Aranthangi, Pudukkottai, Tamilnadu, India.

*Corresponding e-mail: Aarthii.bgk@gmail.com

Abstract – In recent years, video surveillance has become increasingly dependent on automated human tracking via camera networks. Not only is it difficult to follow people over camera networks since people's appearances change over time, but there are also a lot of real-world uses for this technology, including security surveillance, retail, and medical. Due to differences in appearance brought on by changes in viewpoint and illumination, background, occlusion, non-rigid deformations, and intra-class heterogeneity in shape and position, human tracking is one of the most difficult components of the task. In this work a novel deep learning-based model is proposed to overcome these challenges. The input CCTV images are -pre-processed using Wienmed filter and annotated to remove the noise and blurriness in the image. The enhanced images are fed into the YOLO V4 model for tracking the humans in the supermarket with high rate of accuracy. According to the result, the proposed model attains 99.51% of accuracy rate for the detection of human tracking. The modified YOLOV4 yields 99.51% accuracy as opposed to Mask RCNN, YOLO V2, and YOLO V3 are 0.68%, 0.65%, and 1.16%. The CNN, Deep SNORT and YOLO V5 are all outperformed by the proposed model in terms of overall accuracy by 24.51%, 12.51%, and 3.21%, respectively.

Keywords – Human tracking, CCTV images, YOLO V4, Deep learning.

1. INTRODUCTION

Human tracking is a pivotal technology in modern surveillance systems, and its application in supermarkets presents unique opportunities and challenges. The ability to track customers within a retail environment enables store managers to gain insights into shopping behaviours, optimize store layouts, and enhance security [1, 2]. As supermarkets continue to integrate advanced technologies, the use of video surveillance for human tracking becomes increasingly important. In supermarkets, human tracking allows for a detailed analysis of customer movements and interactions with products. This information can be used to improve

customer experience by identifying high-traffic areas, optimizing product placements, and reducing bottlenecks [3, 4]. Additionally, human tracking enhances security by monitoring for suspicious behaviour, thereby helping to prevent theft and ensure the safety of both customers and staff.

Deep learning (DL) [5] has become an effective tool to tackle these issues. DL models may learn to identify and follow people with a high degree of accuracy by utilizing big datasets and intricate neural network designs [6]. The Convolutional neural networks (CNN) [7], in particular, have shown great promise in extracting features from video frames that are crucial for identifying and tracking people in real time. Despite its benefits, human tracking in supermarkets poses several challenges [8]. The dynamic and crowded environment of a supermarket, with its diverse customer base, varying lighting conditions, and frequent occlusions caused by shelves and other customers, complicates the tracking process. These factors make it difficult to maintain accurate and continuous tracking of individuals as they move throughout the store. One major challenge is the requirement for extensive labelled data to train the models effectively. Collecting and annotating such data can be time-consuming and expensive. Additionally, the models must be robust to variations in appearance caused by different clothing, changes in lighting, and partial occlusions [9].

Future advancements in human tracking within supermarkets will likely involve the integration of multiple cameras and sensors, along with improvements in DL algorithms to handle the specific challenges of the retail environment. By addressing these challenges, supermarkets can leverage human tracking technology to enhance customer experience, optimize operations, and improve security [10,16]. In this work a novel deep learning-based

model is proposed to overcome these challenges. The main contribution of the proposed model is as follows.

- In this research, a novel deep learning based model is proposed for the detection of human tracking in supermarket.
- The input CCTV images are -pre-processed using Wienmed filter and annotated to remove the noise and blurriness in the image.
- The enhanced images are fed into the YOLO V4 model for tacking the humans in the supermarket with high rate of accuracy.

The remaining research was arranged in the following manner. The linked publications are fully described in Section 2, the suggested paradigm for tracking people in supermarkets is thoroughly discussed in Section 3, and study findings and views are presented in Section 4. The document is concluded in Section 5 with some further research recommendations.

2. LITERATURE SURVEY

The detection of human tracking forms using deep learning techniques has been part of numerous studies in recent years. A quick summary of a few current research papers is given in the section as follows.

In 2020 Kadim, et al., [11] suggested a single deep-learning-based object tracking system to take use of its exceptional ability to predict objects' appearance even at night. CNN that has already been trained are combined with fully connected layers in the algorithm. The findings indicate that the Adam optimizer, with a sampling ratio of 2:1 and a learning rate of 0.00075 for both positive and negative training data, yields the best accuracy.

In 2020 Yuan, et al., [12] A Simple Baseline for Pose Tracking in Videos of Congested Scenes was proposed. A pipeline for resolving the video's pose tracking issue is presented in this research. Apply DeepSORT to carry out box association, assign ID for the boxes, and make adjustments on two perspectives after getting human positions. Next, use the efficient single-person pose estimation model to provide precise pose predictions. The suggested optical flow smoothing approach then smooths the predictions. Finally, the recommended stance tracking reaches 87%.

In 2022 Rozhbayani, et al., [13] provides a novel deep learning method based on the modified YOLOv5 model for real-time human tracking detection (RHTD) for surveillance footage. to assess the suggested network's performance using two distinct bespoke datasets. For the first dataset, the suggested model achieves 98.3%, and for the second dataset, it is 96.3%. Real-time video frames were used to evaluate the suggested approach, and the results showed very accurate detection.

In 2024 Ji, et al., [14] offer a unique model that uses feature sharing to accomplish the process of spotting abnormal behaviors after human target tracking. This model is based on the tracking of human behaviors. These features are coupled with CNN and LSTM to realize human abnormal behavior recognition. The convolutional layer in MDnet

serves as the input of the abnormal behavior recognition network. The suggested model has 92.1% recognizing rate.

In 2024 Selvakumar, et al., [15] Create a new and effective HABRT model with deep architectures to identify anomalous behavior and lower crime. putting into practice an improved MD-RAN model to identify and categorize anomalous human behavior from real-time surveillance footage. creating an AM-YOLO V3 model to monitor unusual human behavior and avert unnecessary circumstances. to assess how well the created HABRT framework performs in comparison to current classifiers and traditional methods, achieving a 98.55% success rate.

From the above survey, the studies attains several limitations include potential overfitting due to reliance on specific datasets and the challenge of generalizing results across diverse scenarios. The focus on optimizing deep learning models like CNNs and YOLO might neglect computational efficiency, which is critical for real-time applications. Additionally, the models often depend on pre-labeled data and specific conditions (e.g., lighting or crowd density), which may limit their adaptability in dynamic and unpredictable environments. Finally, the high accuracy rates reported might not fully account for real-world variances and practical deployment challenges.

3. PROPOSED METHOD

In this research, a novel DL based model is proposed for the detection of human tracking in supermarket. The input CCTV images are -pre-processed using Wienmed filter and annotated to remove the noise and blurriness in the image. The enhanced images are fed into the YOLO V4 model for tacking the humans in the supermarket with high rate of accuracy.

3.1. Data pre-processing

The Wienmed filter is a sophisticated technique used for deblurring and denoising images, particularly in cases where the image has been degraded by a known blurring function and noise. The primary goal of the Wiener filter is to minimize the mean square error between the estimated image and the original image. A combination of the median and Wiener filters is the Wienmed filter. When combined, these two filters effectively lessen noise dispersion and frame faults. Pre-processing involves improving video frame segments with undesirable distortions as well as a number of image properties that are vital for further processing. The dataset is sent through the Wienmed filter, yielding a pre-processed frame. Equation (1) is utilized to determine mean m in every pixel using gh .

$$F = \frac{1}{rt} \sum_{r,t \in g} h(u, v) \quad (1)$$

Here, r, t is indicates pixel dimensions of every picture, a symbolises every picture and t is an image. Additionally, Gaussian noise variation is provided in Equation (2), σ^2 is variance.

$$\sigma^2 = \frac{1}{rt} \sum_{r,t \in g} h(u, v) - F^2 \quad (2)$$

$$T(r, t) = \sigma^2 [N - h(r, t)] \quad (3)$$

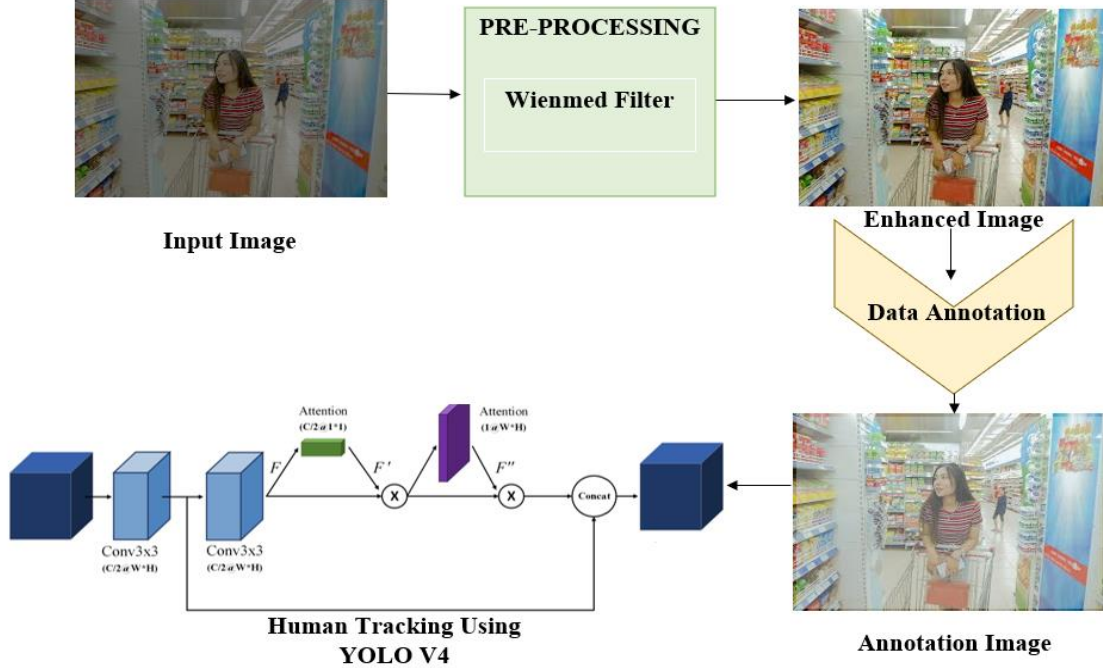


Figure 1. Schematic Illustration of proposed model

Thus, Wienmed filter remove noise in frames and attained images are moved to feature extraction process. The Wiener filter is a powerful tool for deblurring images degraded by known blur and noise. By leveraging the frequency domain and balancing between deblurring and noise suppression, it provides a robust solution for restoring image quality.

3.2. Data Annotation

Data annotation involves labelling images with bounding boxes around relevant objects such as shoppers, shopping carts, shelves, and products. This process is essential for deep learning-based human tracking and activity recognition in supermarkets. To train a machine learning model for human tracking, visual information within photos or videos is labelled and categorized. Annotators identify and label various objects and individuals within the supermarket, including customers, carts, and aisles. Each object or person of interest is marked with bounding boxes indicating their location and dimensions within the image, and annotators also name the type of object or activity. This annotated dataset trains a computer vision model to accurately recognize and classify human activities and objects in the supermarket. This process is crucial for developing a reliable system for tracking customer behavior, optimizing store layouts, and improving overall shopping experience and security.

3.3. Human tracking using YOLO V4

A real-time object recognition model called YOLO V4 has potential uses in many different contexts, one of which is the detection of counterfeits. extremely designed to classify and recognize objects within still or moving frames of vision. To detect signs of visual manipulation for fraud detection,

YOLO V4 may be adjusted. This instantly transforms the raw image into a pattern of cells in order to forecast the boundaries and subclass estimations. YOLO V4 typically uses the robust backbone of a connection, such as CSPDarknet53 or CSPResNeXt50, to extract attributes from the image that is being analyzed. The foundation of neural models consists of multiple layers of segmentation and pooling methods that learn structured properties at different levels of complexity. In the first YOLOv4 backbone, an additional SPP block coupled with CSPDarknet53 connected to the PANet replaced the features of the structure systems used in previous YOLO variants. An effective technique for finding entities of different sizes is used by the SPP.

In PANet, the amount of the featured plane that was initially provided is doubled by upsampling ambiguous feature data fields from the framework and the SPP. By making use of the fundamental YOLO recognition structure, YOLOv4 preserves the head component of YOLOv3 while utilizing an improved backbone system and CSPDarknet53. To enhance the responsive region, it additionally utilizes the notion of spatial system pooling and chooses PANet as the neck component for feature combination. In the YOLOv4 approach, there are three anchoring boxes in each and every cell. Throughout formation, the most important anchor element of the cross-section over association is still in charge of identifying the item being trained, after which it chooses the target from the test set. The following method can be used to anticipate a certain bounding box and its perimeter.

$$y_l = \sigma(z_l) + b_l \quad (4)$$

$$y_m = \sigma(z_m) + b_m \quad (5)$$

The upward variance of the grid is represented by b_l and b_m in the equations above, while the diameter and height are

represented by S_k , S_t , respectively. The center theme of the edge cases is represented by y_l , y_m , while the predicted length and width of the frame are represented by, y_k and y_t , respectively.

$$y_k = \sigma(S_k)e^{z_i} \quad (6)$$

$$y_t = \sigma(S_t)e^{z_j} \quad (7)$$

where the sigmoid rate is represented by σ . and the rate predicted by the model is represented by the variables z_l , z_m , z_i and z_j . YOLOv4 uses these data to carry out the phase modification detection method. In packed object recognition networks, bounding box analysis is a commonly used technique to predict the position of boxes on the visual inputs. The following scale-invariant assessment measure intersection over union was created to evaluate the efficacy of target object identification.

$$B_{IoU} = \frac{B_O \cap B_T}{B_O \cup B_T} \quad (8)$$

Where $B_O \cap B_T$ and $B_O \cup B_T$ are the portions of the detected image bounding box B_T and the estimated box B_O , respectively, that are identified as a collision and a relationship. The object identification objective's loss value (l), which includes YOLO, can be expressed as follows:

$$\Delta_{lf} = \theta_{cr} + \theta_{B_{IoU}} + \theta_{cl} \quad (9)$$

The YOLOv4 is a single-stage, high-precision object recognition system that turns item identification tasks into statistical problems by generating bounding box locations and assigning probabilities to each category. To create Darknet53, the ResNet network framework is connected to the remaining system components. When data from multiple pattern stage inputs are received by the remaining module, higher-level attribute maps are generated. As a result, compared to the ResNet system, network variables are much decreased, feature training capabilities are enhanced, and feature data that remains is improved.

Human Tracking

DeepSORT (Deep Metric Learning for Person Re-identification and Tracking) is a powerful approach for human tracking in videos. It combines the strengths of object detection and data association to provide robust tracking even in challenging scenarios. DeepSORT goes beyond just bounding boxes. It extracts deep features from the detected bounding boxes using a pre-trained deep learning model like YOLO. These features capture the person's appearance, encoding information like clothing color or body shape into a high-dimensional vector. his approach utilizes a classic two-stage multi-object tracking scheme. The first stage involves object detection, where a deep learning model like YOLO or SSD identifies people in each video frame and generates bounding boxes.

Following the detection phase, DeepSORT takes over. It performs the association process, linking bounding boxes across frames to establish unique IDs for each tracked person. Deep features are extracted from bounding boxes using a pre-trained model like YOLO. Cosine similarity between these features and existing tracks is calculated. High similarity suggests the bounding box likely belongs to the

same person. This combined approach allows DeepSORT to handle challenging scenarios like temporary disappearances or similar-looking individuals. The baseline DeepSORT method can be further improved by incorporating a model specifically trained to extract even more informative deep appearance descriptors. This can enhance the accuracy of track association, especially in situations with cluttered backgrounds or similar clothing styles. DeepSORT generates unique IDs for each person detected in the video using deep features and the bounding box information. Individual human poses are estimated for each bounding box identified in the human tracking results. The IDs assigned by DeepSORT are automatically moved to the final poses, associating each pose with the corresponding bounding box and tracked person. This breakdown highlights the role of DeepSORT in human tracking within a CCTV system after bounding boxes are generated by the object detection stage. The additional section on deep appearance descriptors and pose estimation provides further context if applicable to your specific use case.

4. RESULT AND DISCUSSION

In this section, DukeMTMC-VideoReID dataset [16] image is used for tracking humans in supermarket. Although not specifically for supermarkets, this dataset is suitable for person re-identification and tracking in video surveillance. It contains annotated videos with multiple camera views, capturing people moving through a campus environment, which can be analogous to tracking in supermarkets.

4.1. Performance Analysis

This technique is a crucial component of the total face recognition accuracy rate for improving face detection accuracy while reducing the number of false negatives and positives. The simulation results of masked face detection are shown in the Table.1. The parameters used for masked face detection are ACC, PRE, and F1, which are defined as follows:

$$accuracy = \frac{t_p + t_n}{tp + t_n + fp + fn} \quad (10)$$

$$Precision = \frac{t_p}{t_p + f_p} \quad (11)$$

$$F1 = \frac{2(precision * recall)}{precision + recall} \quad (12)$$

where t_p means true positive, and t_n denotes to true negative. False positive (f_p) refers to the false positive, whereas f_n specifies the false negative values of sample images. Recall and precision are important metrics for assessing how well an object recognition model is working. Equation (14) is used to compute Mean Average Precisions (mAP), which is used to assess the YOLO model's objective performance.

$$mAP = \frac{\sum_k \int_0^1 t(i) di}{k} \quad (13)$$

$$fps = \frac{1}{T} \quad (14)$$

An object detection algorithm's frames per second (FPS) indicates how many picture frames it can process in a second.

The calculation of the FPS parameter involves taking the inverse of the computation time expressed in seconds.

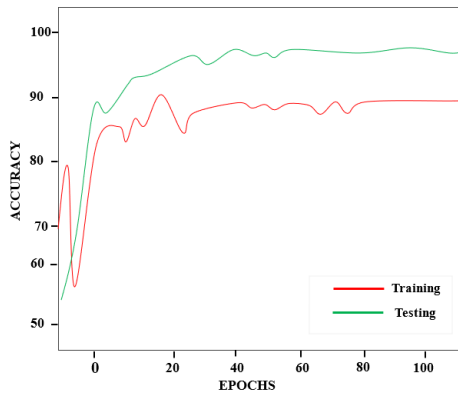


Figure 2. Accuracy of proposed model

In Figure.2, the accuracy graph was generated using 100 epochs and an accuracy range. The accuracy of the YOLOV4 similarly improves as the number of epochs rises.

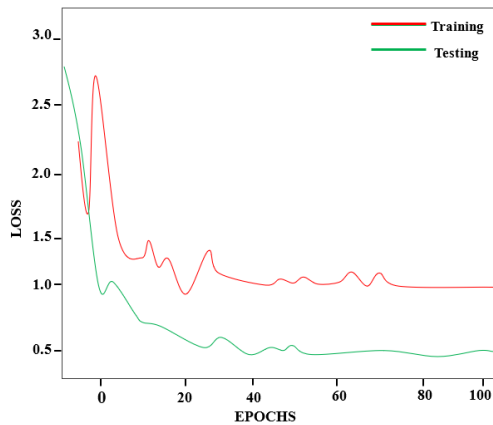


Figure 3. Loss of proposed model

The epochs and loss range are shown in Figure 3, which illustrates that the loss of the proposed model decreases when the epochs are raised. In order to get the optimum evaluation accuracy, the quantity of training epochs was initially calculated in this research. The results demonstrate that after 100 training epochs, the proposed model is detected and attains a low error rate.

4.2. Comparison Analysis

Each neural network's efficacy was evaluated in order to verify that the suggested model generates outcomes with exceptional precision. Competency tests using YOLO V2, YOLO V3, and Mask RCNN deep learning classifiers were performed on the suggested model. The suggested model has a 99.51% accuracy rate, which is greater than typical DL systems.

Table.1 displays the highest detection rate for each of the conventional DL networks after assessing them all. While the accuracy of Mask RCNN, YOLO V2, and YOLO V3 is 2.88%, 1.4%, and 1.26%, the improved YOLOV4 offers 99.51% accuracy.

Table 1. Comparison with traditional model

Networks	Accuracy	Precision	Recall	F1-Score
Mask RCNN	96.63	96.49	97.79	97.12
YOLO V2	98.11	97.58	98.58	97.22
YOLO V3	98.25	98.14	97.66	98.55
Proposed YOLOV4	99.51	99.28	99.38	99.67

Table 2. Contrasts the proposed YOLO forgery models with existing models

AUTHOR	METHODS	ACCURACY
Kadim, Z et al., [11]	CNN	75%
Yuan, L et al., [12]	Deep SNORT	87%
Rozhbayani, G.M et al., [13]	YOLO V3	96.3%
Proposed Model	Modified-YOLOV4	99.51%

In Table 2. A variety of criteria are used to assess the performance of the currently available models with great detection accuracy. CNN, Deep SNORT, and YOLO V3 are all exceeded by the proposed model in terms of overall accuracy by 24.51%, 12.51%, and 3.21%, respectively. When comparing the suggested network to the earlier networks, information was used more quickly.

5. CONCLUSION

In this research, a novel DL based model is proposed for the detection of human tracking in supermarket. The input CCTV images are -pre-processed using wienmed filter and annotated to remove the noise and blurriness in the image. The enhanced images are fed into the YOLO V4 model for tacking the humans in the supermarket with high rate of accuracy. According to the result, the proposed model attains 99.51% of accuracy rate for the tracking human in supermarket. The modified YOLOV4 yields 99.51% accuracy as opposed to Mask RCNN, YOLO V2, and YOLO V3 are 2.88%, 1.4%, and 1.26%. The CNN, Deep SNORT and YOLO V5 are all outperformed by the proposed model in terms of overall accuracy by 24.51%, 12.51%, and 3.21%, respectively. In future, the proposed model is enhanced with advance YOLO models for improving the accuracy performance for tracking humans in supermarket.

CONFLICTS OF INTEREST

This paper has no conflict of interest for publishing.

FUNDING STATEMENT

Not applicable.

ACKNOWLEDGEMENTS

The author would like to express his heartfelt gratitude to the supervisor for his guidance and unwavering support during this research for his guidance and support.

REFERENCES

- [1] X. Zhang, Z. Xu, and H. Liao, "Human motion tracking and 3D motion track detection technology based on visual information features and machine learning", *Neural Computing and Applications*, vol. 34, no. 15, pp. 12439-12451, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

- [2] L. Jiao, D. Wang, Y. Bai, P. Chen, and F. Liu, "Deep learning in visual tracking: A review", *IEEE transactions on neural networks and learning systems*, vol. 34, no. 9, pp. 5497-5516, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [3] S. Gokul Pran, S. Raja, "An efficient feature selection and classification approach for an intrusion detection system using Optimal Neural Network", *Journal of Intelligent & Fuzzy Systems*, vol. 44, no. 5, pp. 8561-71, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [4] J. Pegoraro, D. Solimini, F. Matteo, E. Bashirov, F. Meneghello, and M. Rossi, "Deep learning for accurate indoor human tracking with a mm-wave radar", In *2020 IEEE Radar Conference (RadarConf20)*, pp. 1-6, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [5] S.G. Pran, S. Raja, S. Jeyasudha, "Intrusion detection system based on the beetle swarm optimization and K-RMS clustering algorithm", *International Journal of Adaptive Control and Signal Processing*, vol. 38, no. 5, pp. 1675-89, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [6] Y. Xu, X. Zhou, S. Chen, and F. Li, "Deep learning for multiple object tracking: a survey", *IET Computer Vision*, vol. 13, no. 4, pp. 355-368, 2019. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [7] B. Sivasankari, M. Shunmugathammal, A. Appathurai, and M. Kavitha, "High-Throughput and Power-Efficient Convolutional Neural Network Using One-Pass Processing Elements", *Journal of Circuits, Systems and Computers*, vol. 31, no. 13, pp. 2250226, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [8] J. Fan, W. Xu, Y. Wu, and Y. Gong, "Human tracking using convolutional neural networks", *IEEE transactions on Neural Networks*, vol. 21, no. 10, pp. 1610-1623, 2010. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [9] S.M. Marvasti-Zadeh, L. Cheng, H. Ghanei-Yakhdan, and S. Kasaei, "Deep learning for visual tracking: A comprehensive survey", *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 5, pp. 3943-3968, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [10] G. Zheng, and Y. Xu, "Efficient face detection and tracking in video sequences based on deep learning", *Information Sciences*, vol. 568, pp. 265-285, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [11] Z. Kadim, M.A. Zulkifley, and N. Hamzah, "Deep-learning based single object tracker for night surveillance", *International Journal of Electrical & Computer Engineering (2088-8708)*, vol. 10, no. 4, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [12] L. Yuan, S. Chang, Z. Huang, Y. Zhou, Y. Chen, X. Nie, F.E. Tay, J. Feng, and S. Yan, "A simple baseline for pose tracking in videos of crowd scenes", In *Proceedings of the 28th ACM International Conference on Multimedia*, pp. 4684-4688, 2020. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [13] G.M. Rozhbayani, A. Tuama, and F. Al-Azzo, "Real-Time Human Detection and Tracking Based on Deep Learning Technique", *NeuroQuantology*, vol. 20, no. 6, pp.2084, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [14] X. Ji, S. Zhao, and J. Li, "An algorithm for abnormal behavior recognition based on sharing human target tracking features", *International Journal of Intelligent Robotics and Applications*, pp. 1-13, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [15] S. Selvakumar, "An effective framework of human abnormal behaviour recognition and tracking using multiscale dilated assisted residual attention network", *Expert Systems with Applications*, vol. 247, pp. 123264, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [16] X. Zang, G. Li, W. Gao, and X. Shu, "Exploiting robust unsupervised video person re-identification", *IET Image Processing*, vol. 16, no. 3, pp. 729-741, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

AUTHORS



Aarthi Gopalakrishnan Completed her B. Tech in Information Technology and M. Tech in Computer and Communication from the Anna University in 2011. At present, she is pursuing Ph. D from B.S.Abdur Rahman Crescent Institute of Science and Technology. She has published/presented several research papers in National/ International Conferences.



A. Neelima is Research Scholar at Koneru Lakshmaiah Education Foundation, Department of Computer Science and Engineering, Vaddeswaram, Andhra Pradesh, India. She received B. Tech degree in Aurora's Scientific and Technological Institute, Department of Computer Science and Engineering in 2014, she holds a M. Tech degree in ACE Engineering College, Department of Computer Science and Engineering in 2017. Her research areas are Image Processing, Bioinformatics, Machine Learning and Deep Learning.



C. Priya received her BE degree Computer Science and Engineering from Mepco Schlenk Engineering College, Sivakasi, Tamilnadu, India in 2008 and ME degree in Computer Science and Engineering from Pannai College of Engineering and Technology, Sivagangai, Tamilnadu, India in 2015. She is currently pursuing PhD in the Department of Information and Communication Engineering, Anna University Chennai, Tamilnadu, India. Her areas of interest are Networks and Security, Image processing, Machine Learning and Human Computer Interaction. She is currently serving as a Lecturer in the Department of Computer Engineering at Government Polytechnic College, Aranthangi, Pudukkottai, Tamilnadu.

Arrived: 02.01.2024

Accepted: 04.02.2024