

Analysis of real-time multi-surveillance detection model using YOLO v5

Tapas Pramanik¹, Prakash Gajananrao Burade¹, Sanjeev Sharma²

¹Department of Electrical and Electronics Engineering, Faculty of Engineering, Sandip University, Nashik, India

²Department of Electronic and Communication Engineering, New Horizon College of Engineering, Bangalore, India

Article Info

Article history:

Received May 17, 2024

Revised Dec 26, 2024

Accepted Feb 27, 2025

Keywords:

Optical

Realtime analysis

Surveillance

YOLO

ABSTRACT

Implementation of this advanced nighttime monitoring system provides one of the basic requirements toward the creation of an intelligent urban environment. The nighttime effective monitoring is highly enabled due to seamless integration of multi-directional cameras working as advanced sensors enhancing security measures in smart cities. This paper addresses the mentioned issues directly by proposing the you only look once version 5 (YOLOv5) model dedicated to object detection. It is experimentally confirmed, based on the dataset results, that the mean average precision of YOLOv5 multi-scale (YOLOv5MS) reaches an impressive 88.7%. The results unmistakably confirm domination of the model and its good ability to work over a network of more than 50 security cameras under the high restrictions of our operation. The use of state-of-art nighttime surveillance systems is an important constituent element in the construction of smart urban environment. The smooth interaction between multiple-angle cameras, which work as perceptive sensors, substantially upgrades the functionality of nighttime surveillance and strengthens security practices for smart cities. The current work presented the YOLOv5 model specifically designed for the task of target detection, targeting these issues head-on. The empirical data obtained from the dataset point to an outstanding mean average precision (mAP) of 88.7% for YOLOv5MS. Such results clearly prove the superiority of the model and demonstrate its excellent performance in a network of more than 50 security cameras under our harsh operational conditions.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Corresponding Author:

Sanjeev Sharma

Department of Electronics and Communication Engineering, New Horizon College of Engineering
Bengaluru, India

Email: sanjeevietb@rediffmail.com

1. INTRODUCTION

Video surveillance systems are one of the crucial components of the infrastructure of an intelligent city, motivated by technologies and escalating security issues [1], [2]. The huge volumes of footage that the surveillance systems generate each day pose challenges while also requiring a massive workforce for processing, and even results in inefficient detection of pedestrians at times. Advanced deep learning-based models of detection within smart city frameworks suggest an excellent avenue for improvements in urban security while offering reduced costs and labour usage [3]. These comprise models designed towards continuous recognition of pedestrians, exploiting extensive datasets, which make them viable for use in different environments. This is notably helpful for nighttime surveillance, where identification accuracies tend to be degraded. These solutions enhance the monitoring efficiency and security monitoring by limiting notifications to the security personnel only when a relevant target has been identified.

The role that high-speed communication networks, including 4G and 5G, play in incorporating video surveillance technologies within a coherent monitoring structure for smart cities is of extreme importance. This allows real-time video feeds to be viewed using the intelligent infrastructure of the city; therefore, potential security threats can always be continuously overseen. Installation of such infrastructure is not without challenges. In the experimental smart park environment, 220 cameras were installed, but limitations in hardware and network bandwidth were encountered to analyze all streams simultaneously. In real-time data transmission to the cloud for 50 cameras, a peak bandwidth requirement of 183 Mbps was achieved. The analysis done to scale this for each camera showed that about 805 Mbps was required, making high-speed uplink costs too expensive. Additionally, a large percentage of the data transmitted is redundant, which worsens the strain on available resources. This has led to a significant interest in edge computing solutions as a better approach to handle and process video feeds at the local level [3]. The Nvidia are contemporary edge computing platforms that enable data processing close to the source [4], [5].

Although they have various benefits, the scalability of these systems is limited because they can process only a finite number of video feeds, which is insufficient for the vast camera networks usually used in smart cities. The strategy should involve overcoming these limitations. It provides a technique of how high-performance central processing units (CPUs) and graphics processing units (GPUs) could be integrated in an intranet in a smart city. This architecture has the GPUs take over the computational burdens involved with target detection models, whereas the CPUs are capable enough to collect video feeds by multi-threading capabilities. After this processing phase is complete, the extracted data is transferred to the extranet, allowing the concurrent evaluation of multiple video streams at a very reduced cost in networks. Thereby by creating balance concerning both speed and accuracy as well as cost, an effective urban detection system was delivered and YOLOv5s was designed for a smart city setup to make the model work maximally in identifying pedestrians. Techniques for target detection based on deep learning can be broadly classified under two primary categories. This category provides high accuracy at a compromise in computational speed by generating region proposals before the detection process takes place. Several models, such as R. mask region-based convolutional neural network (R-CNN) [6], R. fast R-CNN [7], and J. faster R-CNN [8], [9], indicate the effectiveness of this method. In the second category, speed of inference is improved by bypassing the candidate frame generation, but using single-stage networks such as single shot multibox detector (SSD) [10], [11] and you only look once (YOLO) [12]–[14]. This work uses YOLO-based algorithms for the real-time recognition of target across multiple cameras [15].

Target detection is challenging particularly in urban areas. The cameras are far from the objects, and hence, the pixel representations are small. Also, the detection is difficult due to the large number of data and frequent occlusions in the smart city videos. In the previous studies, several techniques for enhancing the detection models have been explored [16]. González *et al.* [17] improve the performance for the occluded objects by enlarging the prediction frames, and the generalization is challenging.

Kentaro [18] has exploited Darknet53 framework along with K-means for generation of anchor frames, developing the multimodal attention fusion YOLO model, that has been fine-tuned on a nighttime pedestrian identification benchmark; this method performed well under nocturnal conditions. Recently, Suenderhauf *et al.* [19] proposed the combination of hybrid gaussian model with faster R-CNN to overcome challenges constituted by complex backgrounds. Although successful, this approach has disadvantages: long training time and large model sizes. A ratio-aware approach which adaptively adjusts the aspect ratio of images was proposed by Hosang *et al.* [20]. Problems remain, however, especially when dealing with pedestrians that overlap or are occluded since both false positives and false negatives are introduced. Some of the existing works concentrate on the use of size augmentation to improve performance [21]. Although this improvement favors detection, it simultaneously decreases practicality in the presence of challenges like low processing speeds, an increased model size, and high hardware costs. In the current research, we optimized the YOLOv5 framework to overcome such weaknesses. We modified the backbone architecture of the model so that its size is reduced, thus increasing inference speed without compromising accuracy. We also enhanced the detection performance by incorporating the SE module; hence, the system is well-suited for deployment in smart city applications. This research demonstrates the ability of deep learning frameworks, specifically YOLOv5, to transform urban security infrastructures. We utilize advanced hardware and optimized algorithms in a harmonious balance between cost, precision, and operational efficiency to allow scalable and effective surveillance solutions within intelligent urban environments.

2. METHOD

2.1. YOLO basic principle

All stages, from research design and the method of research technique (algorithms, pseudocode, or whatever), testing, and data gathering are described chronologically. The process of study should be referred to, which would support the consideration as scientific. Below and according to the manuscript with

reference, Figures 1, 2 and Table 1 are located centrally. Figure 2 illustrates the effect of series resistances on the (a) I-V and (b) P-V characteristics.

This research focuses mainly on the detection of pedestrian objects with the intention of enhancing the urban safety of smart cities. While real-time responsiveness is most valued in edge deployments, robustness and mean average precision are very significant, particularly in situations where an operation can be envisioned to continue over a more extended period of time. The YOLO series greatly excels at achieving such requirements. In its conceptual architecture, the YOLO model views the task of object detection as regression and uses a neural network that predicts input images concerning their detection bounding boxes and probabilities corresponding to the target classes. The first variant in the series is YOLOv1. To predict confidence and class scores relevant to the boxes from their position and contents within each cell, the input is up sampled to a 448-by-448-pixel resolution and then divided into a grid of 7×7 cells. The target identification model optimized based on the usage of the GPU is also known as YOLOv5. Other recent variants of the already available YOLOv5 architecture: YOLOv5s, YOLOv5n, YOLOv5l [22], and YOLOv5m [23]. YOLOv5s differs from all the network architectures in the sense of various dimensions with the aim of reducing complexity, and is developed as an edge deployable with real-time object detection ability. Many advantages are enjoyed by YOLOv5s: better precision in detection, robustness, and easier deployment. In particular scenarios, its performance is better than YOLOv5n, YOLOv7, and YOLOv8, and it becomes the best one for conducting full inspections by using extensive camera systems. YOLOv5s can be divided into three main components: The excellence of this performance is produced by good synchronizing in many parts, including back and neck. YOLOv5 offers several advantages, such as adaptive anchor frame calculation and mosaic data augmentation that take place at the beginning [24]. The neck part features combination ability is enhanced through cross stage partial (CSP2) structure and the focus and CSP_X structures within the Backbone component [25]. Spatial pyramidal pooling is used by different kinds of sensory fusion domains. Complete intersection over union (CIoU) loss [26] serves as the bounding box's loss function, and non-maximal suppression (NMS) [27] can be applied in case of overlapping targets. Figure 2 illustrated YOLOv5's structure.

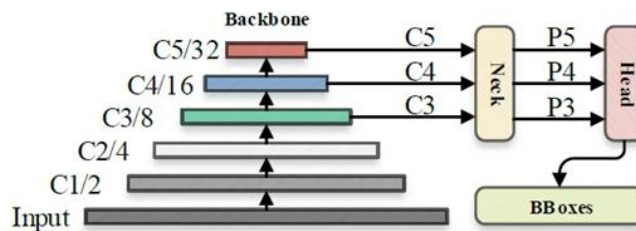


Figure 1. Structure of YOLOv5

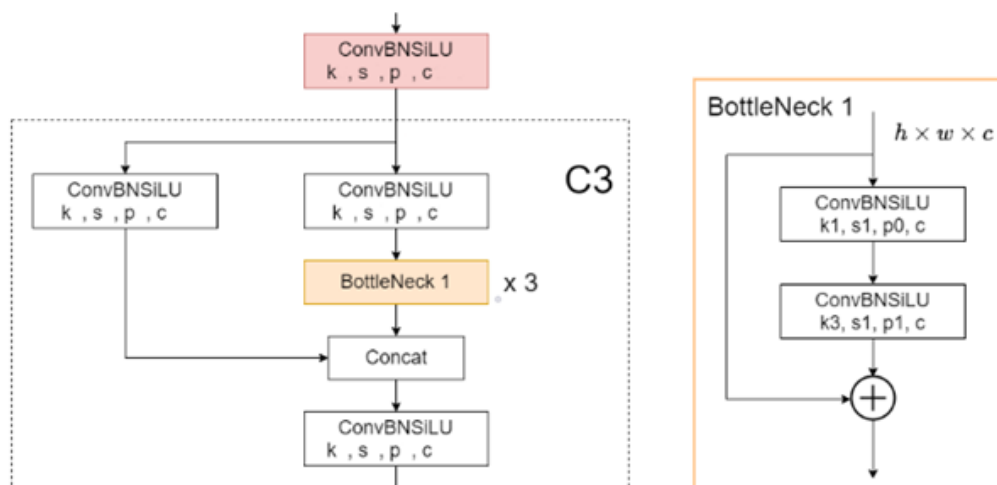


Figure 2. Block structure CSP-Darknet53

2.2. Backbone structure

The structure of YOLOv5 is depicted in Figure 1. The base architecture that is utilized in YOLOv5 is CSP-Darknet53. CSP-Darknet53. YOLO is among the complex neural network architectures using dense and residual blocks to combat the vanishing gradient problem to ensure that information is passed from deeper layers to later layers. However, the use of dense and residual blocks still leads to redundant gradients.

The problem it addresses is cut off from gradients in the flow, whereby, as shown below the YOLOv5 will make use of the methodology developed with CSPNet splitting and interconnecting the feature map partitioning of the base layers with a cross-stage hierarchy. This approach comes with significant advantages to YOLOv5. It not only lowers the count of parameters but also a major number of computations (a few FLOPS, fewer) to increase the inference speed-the important parameter in the model designed to detect real-time objects.

2.3. Neck of YOLOv5

After adding YOLOv5, the neck design of the model experienced two major changes: BottleNeckCSP entered the PANet configuration and SPP entered the odd version of itself.

2.3.1. Path aggregation network

The previous YOLO model is referred to as YOLOv4. The path aggregation network (PANet) structure was used as the feature pyramid network in it. This structure helped mask prediction to be pixel-wise localized and improved the flow of information. The architectural network diagram indicates that the CSPNet method has been implemented in the updates provided for YOLOv5.

2.3.2. Spatial pyramid pooling

This is the output of a certain specified length produced through aggregation of input data, and thus, it provides the benefit of theoretically significant enlargement of the receptive field while distinguishing the most important contextual components without loss to the speed of the network. This block, past versions of YOLO, specifically YOLOv3 and YOLOv4 were employing for the extraction of features from the backbone; while YOLOv5, which entails 6.0 and 6.1, had been using spatial pyramid pooling – fast (SPPF) as just an extension to the SPP block for increasing the speed of the network.

2.3.3. Head of the network

The architecture uses three layers for the tasks of prediction of score, object classification and coordinate detection for the bounding box. Figure 3 depicts the network structure of YOLOv5.

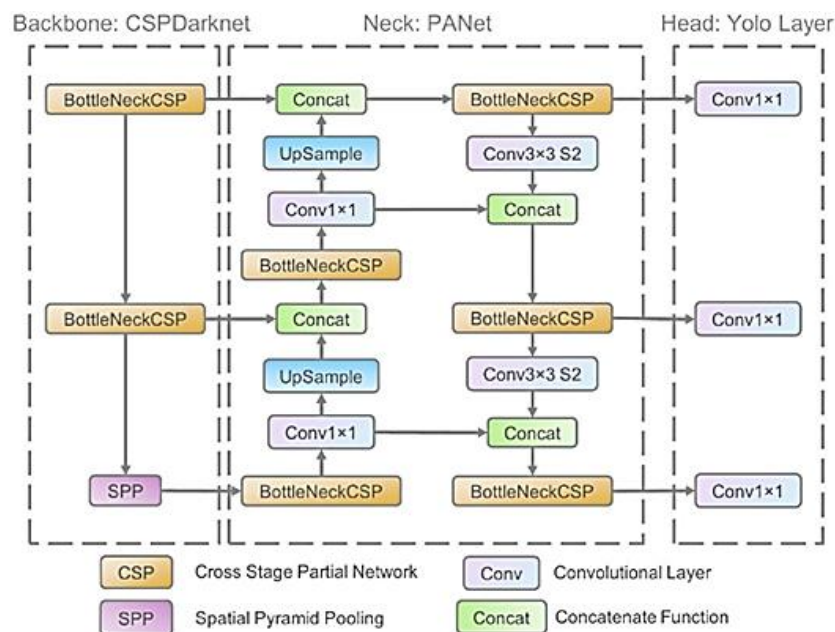


Figure 3. YOLOv5 network structure

2.4. SSAN dataset

Night vision is one of the major factors that affect the performance of human vision, especially at night. Large databases have played a major role in the development of night vision research over time, particularly in the area of picture enhancement. The database used in this experiment contains all the necessary components for enhancing night vision, which classifies a low-light nighttime video as lit if there is minor or significant change in the light. Purchased under a variety of different contexts, night-vision video samples with numerous bright objects featured in the test. Certainly, the major limitation of video samples from night vision camera's is the infrared illumination. The SSAN dataset [15] which comprises nearly 50 video samples for each of all object classes including people and cars has been used in this paper. These are captured at a frame rate of 25 frames per second.

3. RESULTS AND DISCUSSION

This section has analyzed the efficiency of the proposed method. As shown in Figure 4, the YOLOv5 classification is now available for both training and testing. In this regard, samples of different classes can be classified to determine their corresponding classifications, which then allows the trained class to be applied to documents that were not examined before. In order to gauge the effectiveness of the approach taken, some metrics were measured and examined. The performance analysis then compares multiple metrics that are tested for the proposed system against other sub-components of the overall research study. The approach, despite its dependence on critical performance parameters, compares it against various identification algorithms and techniques. For the task of night object detection, the model applied for the classification model used YOLO v5 as it performed with a higher success rate than any method existing at that time.

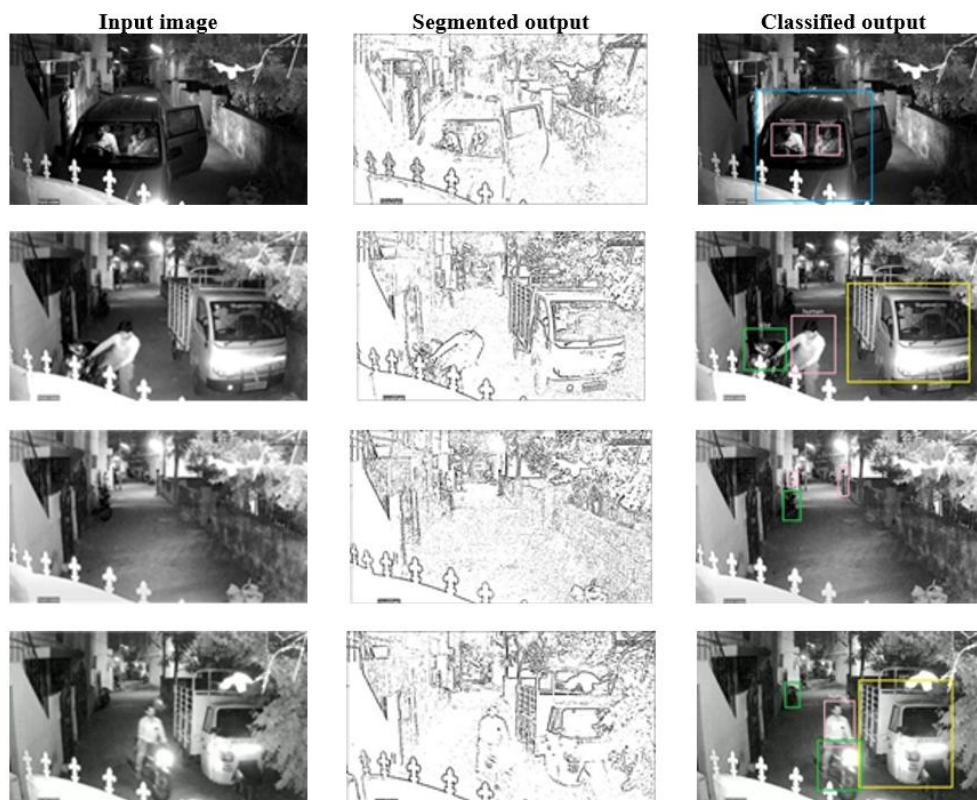


Figure 4. Segmented and classified output of YOLO

It is commonly known as mean average precision, and is widely used in the item identification tasks as a measure to test the efficiency of the proposed model. The data obtained reveals without a shadow of doubt that against other prevailing models currently like faster RCNN and YOLOv3, Figure 5 and Table 1 describes how the proposed model might distinguish objects efficiently and accurately in low light and

vehicle source light. Obviously, the provided YOLOv5 model is supposed to ensure accurate identification results at low light and when objects have a low resolution and noise with bad information. And the architecture provided has shown good efficiency simultaneously without stretching out the inference time - it was able to locate any objects on the street surveillance shot correctly under night conditions.

Table 1. Detection performance is expressed in %, and the detection speed is with multi-scale

Methods	Time	mAP
Faster R-CNN	190.2	65.02
YOLOv3	39.57	87.9
YOLOv5	33.96	88.7

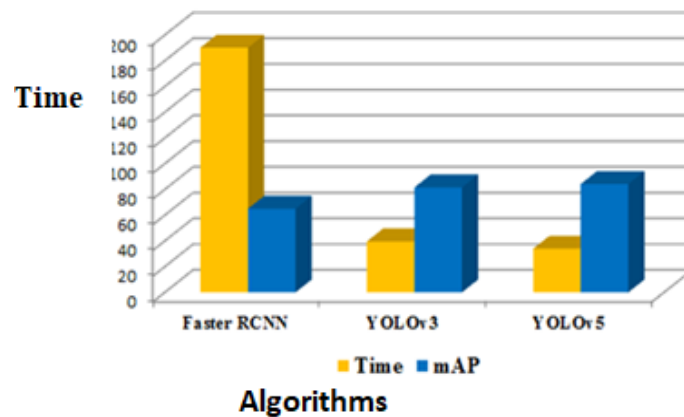


Figure 5. Comparative analysis

4. CONCLUSION

The investigations were based on the widely known outcomes of YOLOv5, which were formulated and tested for detecting and tracking objects in standard nighttime visual imagery within observational contexts. The SSAN dataset was used for experimental analysis. A preliminary study on some of the top-notch detectors like YOLOv3 and faster R-CNN have been performed for the purpose of decide which network could work at best for persons detection at low lighting and source light points emitted by automobiles in street surveillance in nighttime. The performance results yielded by the YOLOv5, at night, faster R-CNN and YOLOv3-control detectors performed similarly. Nonetheless, the YOLOv5's classification was more precise and effective. The average precision of the suggested model was 88.7% in 33.96 milliseconds, which corresponds well with a good average accuracy. The proposed model was yet capable of detecting source light points that originated from vehicular backgrounds and objects in low illumination scenes. Thus, this is well-grounded for further developments of this model.

ACKNOWLEDGEMENTS

Author thanks to the Sandip University for giving facility for the research work.

FUNDING INFORMATION

There is no funding.

AUTHOR CONTRIBUTIONS STATEMENT

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Tapas Pramanik	✓	✓			✓			✓	✓	✓			✓	
Prakash Gajananrao Burade		✓				✓		✓	✓	✓	✓	✓		
Sanjeev Sharma	✓		✓	✓			✓			✓	✓		✓	✓

C : Conceptualization	I : Investigation	Vi : Visualization
M : Methodology	R : Resources	Su : Supervision
So : Software	D : Data Curation	P : Project administration
Va : Validation	O : Writing - Original Draft	Fu : Funding acquisition
Fo : Formal analysis	E : Writing - Review & Editing	

CONFLICT OF INTEREST STATEMENT

There is no conflict of interest.

DATA AVAILABILITY

Not Applicable.

REFERENCES

- [1] F. C. Li, A. Gupta, E. Sanocki, L. He, and Y. Rui, "Browsing digital video," in *Proceedings of the SIGCHI conference on Human Factors in Computing Systems*, Apr. 2000, pp. 169–176, doi: 10.1145/332040.332425.
- [2] Y.-J. Liu, C.-X. Ma, Q. Fu, X. Fu, S.-F. Qin, and L. Xie, "A sketch-based approach for interactive organization of video clips," *ACM Transactions on Multimedia Computing, Communications, and Applications*, vol. 11, no. 1, pp. 1–21, Aug. 2014, doi: 10.1145/2645643.
- [3] G. Jawaheer, P. Weller, and P. Kostkova, "Modeling user preferences in recommender systems: a classification framework for explicit and implicit user feedback," *ACM Transactions on Interactive Intelligent Systems*, vol. 4, no. 2, pp. 1–26, Jul. 2014.
- [4] X. Zhang and G. W. Furnas, "MCVEs: Using cross-scale collaboration to support user interaction with multiscale structures," *Presence: Teleoperators and Virtual Environments*, vol. 14, no. 1, pp. 31–46, 2005, doi: 10.1162/1054746053890288.
- [5] J. F. McCarthy and E. S. Meidel, "ActiveMap: a visualization tool for location awareness to support informal interactions," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 1707, 1999, pp. 158–170.
- [6] K. V. Nesbitt, "Getting to more abstract places using the metro map metaphor," in *Proceedings. Eighth International Conference on Information Visualisation, 2004. IV 2004.*, 2004, vol. 8, pp. 488–493, doi: 10.1109/IV.2004.1320189.
- [7] Y. Hu, E. R. Gansner, and S. Kobourov, "Visualizing graphs and clusters as maps," *IEEE Computer Graphics and Applications*, vol. 30, no. 6, pp. 54–66, 2010, doi: 10.1109/MCG.2010.101.
- [8] W. Ullah, T. Hussain, and S. W. Baik, "Vision transformer attention with multi-reservoir echo state network for anomaly recognition," *Information Processing and Management*, vol. 60, no. 3, 2023, doi: 10.1016/j.ipm.2023.103289.
- [9] C. Szegedy, V. Vanhoucke, S. Ioffe, J. Shlens, and Z. Wojna, "Rethinking the inception architecture for computer vision," in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, vol. 2016-Decem, pp. 2818–2826, doi: 10.1109/CVPR.2016.308.
- [10] W. Ullah, F. U. Min Ullah, Z. Ahmad Khan, and S. Wook Baik, "Sequential attention mechanism for weakly supervised video anomaly detection," *Expert Systems with Applications*, vol. 230, p. 120599, Nov. 2023, doi: 10.1016/j.eswa.2023.120599.
- [11] J. L. Salazar González, J. A. Álvarez-García, F. J. Rendón-Segador, and F. Carrara, "Conditioned cooperative training for semi-supervised weapon detection," *Neural Networks*, vol. 167, pp. 489–501, 2023, doi: 10.1016/j.neunet.2023.08.043.
- [12] I. Jegham, A. Ben Khalifa, I. Alouani, and M. A. Mahjoub, "Vision-based human action recognition: an overview and real world challenges," *Forensic Science International: Digital Investigation*, vol. 32, p. 200901, Mar. 2020, doi: 10.1016/j.fsidi.2019.200901.
- [13] L. Jiao, H. Wu, R. Bie, A. Umek, and A. Kos, "Towards real-time multi-sensor golf swing classification using deep CNNs," *Journal of Database Management*, vol. 29, no. 3, pp. 17–42, 2018, doi: 10.4018/JDM.2018070102.
- [14] Z. Aziz, N. Bhatti, H. Mahmood, and M. Zia, "Video anomaly detection and localization based on appearance and motion models," *Multimedia Tools and Applications*, vol. 80, no. 17, pp. 25875–25895, 2021, doi: 10.1007/s11042-021-10921-0.
- [15] G. Wang, H. Ding, M. Duan, Y. Pu, Z. Yang, and H. Li, "Fighting against terrorism: a real-time CCTV autonomous weapons detection based on improved YOLO v4," *Digital Signal Processing*, vol. 132, p. 103790, Jan. 2023, doi: 10.1016/j.dsp.2022.103790.
- [16] M. Kumar, D. K. Yadav, S. Ray, and R. Tanwar, "Handling illumination variation for motion detection in video through intelligent method: An application for smart surveillance system," *Multimedia Tools and Applications*, vol. 83, no. 10, pp. 29139–29157, 2024, doi: 10.1007/s11042-023-16595-0.
- [17] J. L. S. González, C. Zaccaro, J. A. Álvarez-García, L. M. S. Morillo, and F. S. Caparrini, "Real-time gun detection in CCTV: an open problem," *Neural Networks*, vol. 132, pp. 297–308, 2020, doi: 10.1016/j.neunet.2020.09.013.
- [18] Y. Kentaro, "An experimental study of the effects of landmark discovery and retrieval on visual place recognition across seasons," *Workshop on Visual Place Recognition in Changing Environments at the IEEE International Conference on Robotics and Automation*, 2015.
- [19] N. Suenderhauf et al., "Place recognition with ConvNet Landmarks: viewpoint-robust, condition-robust, training-free," in *Robotics: Science and Systems XI*, Jul. 2015, vol. 11, doi: 10.15607/RSS.2015.XI.022.
- [20] J. Hosang, R. Benenson, P. Dollar, and B. Schiele, "What makes for effective detection proposals?," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 38, no. 4, pp. 814–830, 2016, doi: 10.1109/TPAMI.2015.2465908.
- [21] M. M. Cheng, Z. Zhang, W. Y. Lin, and P. Torr, "BING: binarized normed gradients for objectness estimation at 300fps," in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, 2014, pp. 3286–3293, doi: 10.1109/CVPR.2014.414.
- [22] C. L. Zitnick and P. Dollár, "Edge boxes: locating object proposals from edges," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 8693 LNCS, no. PART 5, 2014, pp. 391–405.

- [23] P. Krähenbühl and V. Koltun, "Geodesic object proposals," in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, vol. 8693 LNCS, no. PART 5, 2014, pp. 725–739.
- [24] P. Arbelaez, J. Pont-Tuset, J. Barron, F. Marques, and J. Malik, "Multiscale combinatorial grouping," in *2014 IEEE Conference on Computer Vision and Pattern Recognition*, Jun. 2014, pp. 328–335, doi: 10.1109/CVPR.2014.49.
- [25] B. Alexe, T. Deselaers, and V. Ferrari, "Measuring the objectness of image windows," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 34, no. 11, pp. 2189–2202, 2012, doi: 10.1109/TPAMI.2012.28.
- [26] E. Rahtu, J. Kannala, and M. Blaschko, "Learning a category independent object detection cascade," in *2011 International Conference on Computer Vision (ICCV)*, Nov. 2011, pp. 1052–1059, doi: 10.1109/ICCV.2011.6126351.
- [27] S. Manen, M. Guillaumin, and L. Van Gool, "Prime object proposals with randomized prim's algorithm," in *2013 IEEE International Conference on Computer Vision*, Dec. 2013, pp. 2536–2543, doi: 10.1109/ICCV.2013.315.

BIOGRAPHIES OF AUTHORS



Tapas Pramanik    received his B. Tech degree in electronics engineering from RTMNU, Nagpur in 2008. He has done PG in M. Teh (VWI) and MBA (marketing and HR) in 2012 and 2020 respectively. He has 14 years of teaching experience. He has published 1 paper in international journal in 2020 and 2 papers in international conferences. He is currently working in Tulsiramji Gaikwad Patil College of Engineering and Technology Nagpur. He can be contacted at email: tapasprmnk86@gmail.com.



Dr. Prakash Gajananrao Burade    received his Ph.D. from RTMNU, Nagpur University in 2012. He is currently serving as a professor and the Dean of the School of Engineering and Technology, as well as an associate with the Electrical and Electronics Engineering Department at Sandip University, Nashik, Maharashtra, India. He has published 55 research papers in prestigious international journals and over 40 papers in international conferences. He holds 5 patents to his name. His research interests include custom power devices, power electronics, power system optimization, and FACTS devices. He can be contacted at email: prakash.burade@sandipuniversity.edu.in.



Dr. Sanjeev Sharma    has received his B.Tech., degree in electronics and communication engineering from M.J.P. Rohikhand University, Uttar Pradesh in 2001. He has done PG, Diploma in VLSI from CDAC. He received M.E and Ph.D. degree in VLSI from Rashtrasant Tukadoji Maharaj Nagpur University, Nagpur in 2008 and 2016 respectively. After 7 years of industrial experience and 15 year of academic experience he joined New Horizon College of Engineering, Bangalore in 2018. Currently he is heading the Department of Electronics and Communication Engineering. His contribution is counted in the field of Low Power VLSI Applications. He has successfully advised 32 masters graduate and currently guiding Ph.D. students under Visvesvaraya Technological University (VTU), Karnataka. He has published 62 technical papers in reputed journals and conference proceedings with 4 international book chapters, authored books, and provisionally filed 39 patents. He has also received best paper awards in ICSESD at Thailand and IJETEMR journal. He also won best Teacher award in the year 2018 from STAMI. He is possessing many professional body memberships like IETE, ISTE, UACEE, and IACSIT. He is editorial board members of IJATES Journal. American Journal of Science, Engineering and Technology and reviewer of Global Research and Development Journal (GRD journal). He is an eminent reviewer for Ph.D. thesis for Nagpur University and Bhagwant University. He can be contacted at email: sanjeevietb@rediffmail.com.