

Characteristics of Action Recognition using graph edge CNN

¹Kalaiaarasan.R,²Sharan.M, ²Sabareeswaran.P

¹Assistant Professor, ²UG Student

Department of Electronics and Communication Engineering,
M.Kumarasamy College of Engineering, Karur, Tamilnadu

Abstract

Body joints, legitimately acquired from a posture estimation model, have demonstrated successful for activity acknowledgment. Existing works center around breaking down the elements of human joints. Notwithstanding, with the exception of joints, people additionally investigate movements of appendages for getting activities. Given this perception, we research the elements of human appendages for skeleton based activity acknowledgment. In particular, we speak to an edge in a diagram of a human skeleton by coordinating its spatial neighboring edges (for encoding the collaboration between various appendages) and its worldly neighboring edges (for accomplishing the consistency of developments in an activity). In light of this new edge portrayal, we devise a diagram edge convolutional neural system (CNN). Thinking about the complementarity between diagram hub convolution what's more, edge convolution, we further develop two cross breed arranges by presenting diverse shared transitional layers to coordinate diagram hub and edge CNNs. Our commitments are two fold, diagram edge convolution and half breed systems for coordinating the proposed edge convolution and the customary hub convolution. Test results on the Kinetics and NTURGB+D informational collections exhibit that our diagram edge convolution is convincing to get the activity qualities and in our graph edge CNN fundamentally beats the current situation to the workmanship skeleton based activity acknowledgment techniques.

Keywords: Convolutional Neural Network, Action Recognition, Skeleton data

I. INTRODUCTION

HUMAN conduct investigation is an essential and testing issue in PC vision. As of late, there develop various errands to handle this issue from unmistakable viewpoints, counting human posture estimation to distinguish and confine major joints of the human body[1],[2], step acknowledgment to distinguish individuals' strolling design [3]–[6], and human activity acknowledgment to order human activities [7]–[9]. We center around human activity acknowledgment thinking about its wide applications in video observation, augmented reality, human–PC connection, and mechanical technology. Old style monocular RGB video-based activity acknowledgment frequently experiences issues in exhaustively speaking to activities in the 3D space [10], [11]. Given the quick advancement of minimal effort gadgets to catch 3D information (example camera exhibits what's more, Kinect), an expanding number of studies are effectively being directed on 3D activity acknowledgment.as multi leveled equations, graphics, and tables are not prescribed, although the various table text styles are provided. The formatter will need to create these components, incorporating the applicable criteria that follow. Skeletons produced from profundity maps are regularly invariant to perspective or appearance. In addition, as an elevated level deliberation of human activities, skeletal information incredibly streamline the trouble in speaking to and understanding diverse activity classifications. As of late, various advancements have been created to evaluate the worldly movement of a skeleton by

following and examining the movement of human joints. One of the most clear methodologies is to connect directions of joints into a long component vector at each time step and afterward input it into fleeting investigation models [15], [16]. In any case, the spatial connection between the joints, as a crucial some portion of a human skeleton, has been disregarded.

To exploit, the associations between the joints, [17] utilized a covariance framework for skeleton joint areas after some time as a discriminative descriptor, and [15] proposed speaking to an activity by a weighted summation of action lets. Most as of late, some profound neural system based arrangements are created. For instance, [18] utilized different repetitive neural systems (RNN) in a tree like chain of command to classify activity classes. Zhang et al. [19] planned a view versatile RNN with long present moment memory (LSTM) design that can adjust to generally appropriate perception perspectives from start to finish. Yan et al. [20] built up the spatial transient chart convolutional systems are used to make an complete analysis of human skeleton on over the spatial and fleeting areas of body. Lee et al.[21] proposed to the troupe worldly of different sliding LSTM systems made out of short, medium, and long haul TS-LSTM to catch different fleeting highlights. Li et al. [22] presented a two stream convolutional neural organize (CNN) to process both crude arranges and movement information got by subtracting joint facilitates in the back to back outlines. Ke *et al.* [23] changed a skeleton grouping into three clasps comparing to three tube shaped directions of the skeleton arrangement before applying deep CNN.

Accepting human joints as diagram hubs, skeleton information can be normally dealt with by diagram CNNs [12], [20]. As of late, there is expanding enthusiasm for broadening profound learning for diagram information. By and large, existing diagram CNNs can be partitioned into two classifications: spatial diagram convolutional arrange (GCN) [24],[25] and otherworldly GCN [26], [27]. GCNs structured in the spatial space mirror the picture based convolution and perform convolution by incorporating every neighboring hub at each position. Otherworldly GCNs first change charts into their range and afterward direct convolution with convolutional piece planned in the ghostly area as indicated by the convolution hypothesis. These current works frequently translate the diagram from the point of view of diagram hub, while dismissing the significance of chart edges, particularly when diagram edge have prominent physical implications in some genuine situation, e.g., appendages in the human skeleton.

Considering this task of skeleton based activity acknowledgment, these techniques have yielded great execution enhancements on certifiable activity acknowledgment informational collections. They generally concentrate on human joints, which are straightforwardly gotten from a present estimation calculation. Notwithstanding, other than joints, people likewise investigate the movements of appendages for getting activities. Subsequently, it is additionally enlightening to focus on the elements of human appendages (relating to the edges in a skeleton diagram) for skeleton based activity acknowledgment.

In view of this perception, in this article, we propose a novel diagram edge convolution activity on human appendages. Considering a edge convolutional in a weighted coordination of neighboring edges are spoken to solitary diagram (see Fig. 1). By coordinating the spatial and transient nearby edges will speak to diagram edge convolutional of transient arrangement. Diagram edge point view are highlighted from edge convolutional systems. Not the same as ordinary diagram convolutional models focusing on chart hubs [20], [28], [29], we center around precious data conveyed by diagram edges, with chart hubs demonstrating just the associations between the edges. Along these lines, we are capable to catch the relationship and conditions between human appendages, which relate to edges in human skeleton charts. Moreover, considering the complementarity between the chart hub and edge convolution, we further plan two half and half systems with various shared moderate layers to

incorporate chart hub and edge-convolutional systems: one is with an extra mutual completely associated layer legitimately blending the highlights from two systems by direct mix, while the other is furnished with two extra shared convolutional layers that remember the two hubs and edges for every convolution. We preference skeleton based human activity acknowledgment for obtaining adequacy in type of graphs.

II. LITERATURE SURVEY

Lin Sun, Kui Jia *et al.*[30] proposed that human activities caught in video groupings are three dimensional signs portraying visual appearance and movement elements. To learn activity designs, existing methods embrace Convolutional or potentially Recurrent Neural Networks (CNNs and RNN). CNN based strategies are powerful in learning spatial appearances, yet are constrained in displaying long haul movement dynamics. However, innocently applying RNNs to video sequences in a convolutional way certainly expect that movements in recordings are stationary crosswise over various spatial locations. This supposition that is substantial for transient movements yet invalid when the length of the movement is long. Not at all like conventional two stream designs which use RGB and optical stream data as information, our two-stream model use the two modalities to together train both info entryways and both overlook doors in the system as opposed to regarding the two streams as isolated elements with no data about the other. We apply this start to finish framework to benchmark datasets (UCF-101 and HMDB-51) of human activity recognition. Investigations show that on both datasets, our genius presented strategy outflanks every single existing one that depend on LSTM as well as CNNs of comparable model complexities

Junwu Weng, Mengyuan Liu[31] proposed that the representation of 3D present assumes a basic job for 3D activity and motion acknowledgment. As opposed to speaking to a 3D present directly by its joint areas, in this paper, we propose a Deformable Pose Traversal Convolution Network that applies one-dimensional convolution to cross the 3D present for its portrayal. Rather than fixing the open field when performing traversal convolution, it upgrades the convolution piece for each joint, by thinking about relevant joints with different loads. This deformable convolution better uses the contextual joints for activity and motion acknowledgment and is increasingly strong to uproarious joints. Besides, by sustaining the scholarly posture highlight to a LSTM, we per-structure start to finish preparing that together upgrades 3D present portrayal and transient grouping acknowledgment. Trials on three benchmark datasets approve the focused exhibition of our proposed strategy, just as its proficiency and heartiness to deal with boisterous joints of posture.

Xin Chen, Jian Weng *et al.* [32] proposed that the Adapting profound portrayals have been applied in real life acknowledgment broadly. Be that as it may, there have been a couple of examinations on the most proficient method to use the basic complex data among various activity recordings to upgrade the acknowledgment precision and productivity. In this paper, we propose to consolidate the complex of preparing tests into profound realizing, which is characterized as profound complex learning (DML). The proposed DML structure can be adjusted to most existing profound systems to adapt progressively discriminative highlights for activity acknowledgment. At the point when applied to a convolutional neural system, DML inserts the past convolutional layer's complex into the following convolutional layer; in this manner, the discriminative limit of the following layer can be advanced. We likewise apply the DML on a confined Boltzmann machine, which can reduce the overfitting issue. Exploratory outcomes on four standard activity databases (i.e., UCF101, HMDB51, KTH, and UCF sports) show that the proposed technique outflanks the cutting edge strategies.

Wenbo Li, Longyin Wen *et al.*[33] presents the RNN Tree (RNN-T), a versatile learning structure for skeleton based human activity acknowledgment. Our strategy sorts activity classes and uses various

Recurrent Neural Networks (RNNs) in a treelike chain of command. The RNNs in RNN-T are co-prepared with the activity class chain of importance, which decides the structure of RNN-T. Activities in skeletal portrayals are perceived by means of a progressive surmising process, during which individual RNNs separate better grained activity classes with expanding certainty. Surmising in RNN-T closes when any RNN in the tree perceives the activity with high certainty, or a leaf hub is come to. RNN-T adequately addresses two principle difficulties of enormous scale activity acknowledgment: (I) ready to separate fine-grained activity classes that are recalcitrant utilizing a solitary system, and (ii) versatile to new activity classes by expanding a current model. We show the adequacy of RNN-T/ACH technique and contrast it and the best in class strategies on an enormous scale dataset and a few existing benchmarks.

III. PROPOSED SYSTEM

By seeing the diagram convolution activity as an exceptional instance of the chart consideration activity, we infer our EGNN(C) layer from the equation of EGNN(A) layer. In reality, the basic distinction among GCN and GAT is whether we utilize the consideration coefficients (i.e., framework α) or the contiguousness lattice to total hub highlights. With this view, we determine EGNN(C) by supplanting the consideration coefficient lattices $\alpha \bullet p$ with the relating edge highlight frameworks $E \bullet p$. The subsequent recipe for EGNN(C) is given as follows:

$$X^l = \sigma \left[\bigcup_{p=1}^P (E_{\cdot p} X^{l-1} W^l) \right],$$

Numerous diagrams are coordinated. As a rule, edge course contains significant data about the diagram. For instance, in a reference organize, AI papers once in a while refer to science papers or other hypothetical papers. Be that as it may, arithmetic papers may only sometimes refer to AI papers. In numerous past examinations including GCNs and GAT, edge bearings are not considered. In their tests, coordinated diagrams, for example, reference systems are just treated as undirected charts. In this paper, we appear in the trial part that disposing of edge headings will lose significant data. By survey bearings of edges as a sort of edge highlights.

In this way, each coordinated channel is increased to three channels. Note that the three channels characterize three sorts of neighborhoods: forward, in reverse, and undirected. Therefore, EGNN will total hub data from these three distinct sorts of neighborhoods, which contains the heading data. Taking the reference organize for example, EGNN will apply the consideration component or convolution activity on the papers that a particular paper referred to, the papers referred to this paper, and the association of the previous two. With such edge highlights, discriminative examples in different kinds of neighborhoods can be adequately captured.

We first audit the old style convolution activity led on the pictures. At that point, we move to our chart edge convolution and GECNNs just as their application to skeleton-based activity acknowledgment. At long last, two half and half models coordinating edge and hub convolutional systems are presented.

IV. GRAPH EDGE CONVOLUTIONAL NEURAL NETWORKS

In this area, we first audit the traditional convolution activity led on the pictures. At that point, we move to our graph edge convolution and GECNNs just as their application to skeleton based activity acknowledgment. At last, two half and half models incorporating edge and hub convolutional systems are presented.

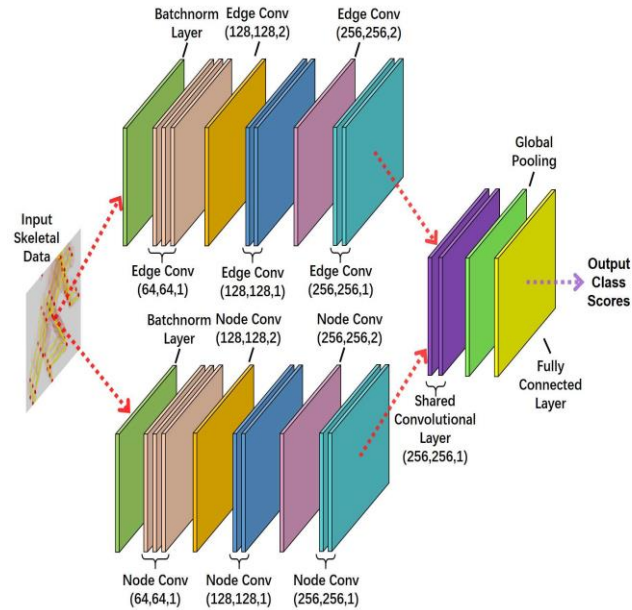


Fig.1. Network structure of the SLHM.

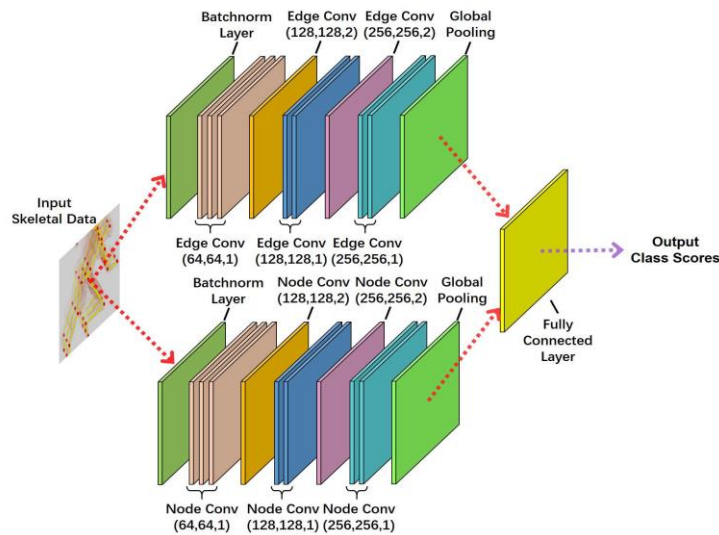


Fig.2. Illustration of Proposed System Network

A. Classical Image Convolution

To make our graph edge convolution progressively reasonable furthermore, show its association with picture convolution, we first present a reformulation of old style convolutions on image pixels. Given a grayscale image, we focus the convolution bit at pixel x_{ij} , where i and j mean the line file and segment file, separately. The convolution yield can be inferred as follows:

$$x_{ij}^{\text{out}} = \sum_{x_{mn} \in N(x_{ij})} x_{mn} w(l(x_{mn}))$$

where x_{ij}^{out} is the yield of this convolution activity, $N(x_{ij})$ speaks to the arrangement of neighboring pixels of x_{ij} , and x_{mn} signifies the pixels in $N(x_{ij})$. By and by, for a grayscale picture,

$N(x_{ij})$ is made out of eight neighboring pixels x_{ij} just as x_{ij} itself, which is the genuine response to the convolution field. l is a marking capacity to allocate a request to every person component in $N(x_{ij})$. Let us consider a 3×3 convolutional of grayscale picture, will deal out $\{1, 2, \dots, 9\}$ to nine pixels in $N(\cdot)$ changes from left to right. The weight work w obtains a pixel weight $w(l(\cdot))$ as indicated by demand. Three various diverts in a RGB picture, gives estimation on every pixel in 3D highlight vector it relating weight turns into a 3D weight vector also. The augmentation between pixels also, loads is stretched out to the inward item.

CNNs have accomplished amazing exhibitions in different situations, e.g., image order and object identification. Be that as it may, it indirect to put on a pixel convolution to chart information, since diagram information don't have the gridded exhibit construction that pictures, movie, sign information do. Every pixel in gridded pictures will have a similar number of neighbors and similar associations regarding a neighbor in guaranteed course. Non gridded graphs don't have these restrictions. A different gridded graph will vary in the neighbors quantity from edge to hub edge and there isn't really a geometrical translation in some random association between two edges.

B. Graph Edge convolution

We speak to a chart by $G = (V, E)$, in which $V = \{v_i = 1, \dots, N\}$ is a hub set with N components and $E = \{e_{ij} | v_i, v_j \text{ are associated}\}$ is an edge set containing all edges of G . A few ideas must be acquainted with encourage the clarification of diagram convolution.

- 1) Path Between Two Edges: A way is a lot of particular hubs and edges interfacing two edges in a diagram. More than one way may exist between two edges.
- 2) Length of Path Between Two Edges: Given a couple of edges, various ways contain various quantities of hubs and where the edges will exist. In such a way the hubs will be characterized as per its length with various lengths.
- 3) Shortest Path Between Two Edges: Among couple of edges these two edges are associating with littlest length characterized with various lengths.
- 4) Separation Between Two Edges: Given a couple of edges, the length of the most limited way is characterized as the separation between the two edges, indicated as D .

C. Convolutional Neural Network:

Information contribution are standardized in group standardization layer. Then focus directions of appendages are formed. At that point, the information are bolstered into a nine-layer GECNN. The initial three layers have 64 channels for yield, the following three layers have 128 channels for yield, and the last three layers have 256 channels for yield. Pooling is to process on fourth and seventh layers. Overall pooling layer is energized in last yield tensor from convolutional layers and for each video progression it gets 256-D feature vector. Before the worldwide pooling layer, the yield tensor is completely associated to acquire the class notches for every video succession.

D. Combining Graph Edge and Node Convolutions

From alternate points of view the edge and node convolutions are set to separate highlights. Both edge convolution what's more, hub convolution have their special favorable circumstances what's more,

accordingly are integral to one another. As we referenced prior, a hub convolutional influences the elements of joints and edge convolutional use the element of appendages. Therefore, we consider to structure half and half models coordinating both two systems to use the benefits of both edge and hub convolution. In particular, we have two half and half systems, counting one utilizing a common completely associated layer and another one furnished with two shared convolutional layers.

1) Sequence Level Hybrid Model: As referenced before, the initial segment of the SLHM includes two isolated streams, counting an edge convolutional stream and a hub convolutional stream. The edge convolutional stream incorporates a standardization layer, trailed by nine edge convolutional layers and a worldwide pooling layer. The standardization, nine hub convolutional, a worldwide pooling layer have same construction like hub convolutional stream. A 256D highlight vector are developed after the pooling layer initial segment. In SLHM, the two 256D yielded tensor will link to solitary tensor after pooling layer and afterward the last characterization result was completely associated to acquire input, which is the class notches for each succession, showing the probabilities of ordering a succession (see Fig. 1).

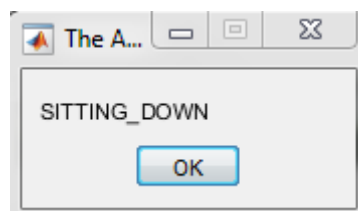
2) Body Part Level Hybrid Model: The are two sections in body part level hybrid model . The initial part is similar like the initial part of the SLHM, then again, actually the pooling layer is displaced. Four layers are incorporates in second part of the body part hybrid model. In the subsequent portion, the two layers of shared convolutional layers are contributed to yield tensors of two previous streams. At that point, worldwide pooling layer is fixed on the yield tensors of the mutual convolutional arrange. At long last, we feed the outcome into a completely associated layer to get the last order result (see Fig. 2).

V. EXPERIMENTAL RESULTS

1) Input image



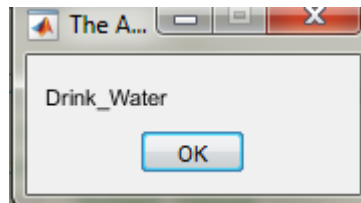
Output result



2) *Input image*



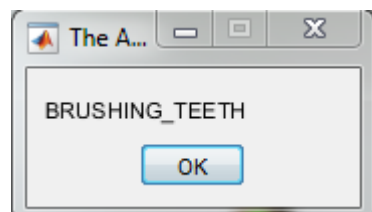
Output result



3) *Input image*



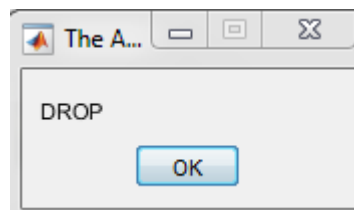
Output result



4) Input image



Output result



CONCLUSION

Thinking about the absence of consideration on human appendage elements for skeleton based activity acknowledgment, we recommend graph edge convolution that speaks to each edge by incorporating its neighboring edges. By leading convolutional graph edges of human skeleton our model catches the relationship and conditions on human appendages. Our model is to bring about surprisingly prevalent execution over the current best in class strategies. Considering the complementarity of node and edge convolutional systems, we further look to consolidate hub convolution and edge convolution in two crossover models that acquire the benefits of both the models. The investigation results show that our half breed models can additionally improve the exhibition.

REFERENCES

1. J. J. Tompson, A. Jain, Y. LeCun, and C. Bregler, "Joint training of a convolutional network and a graphical model for human pose estimation," in *Proc. Adv. Neural Inf. Process. Syst.*, 2014, pp. 1799–1807.
2. Z. Cao, T. Simon, S.-E. Wei, and Y. Sheikh, "Realtime multi-person 2D pose estimation using part affinity fields," 2016, *arXiv:1611.08050*. [Online]. Available: <https://arxiv.org/abs/1611.08050>
3. S.Palanivel Rajan, et.al., "Performance Evaluation of Mobile Phone Radiation Minimization through Characteristic Impedance Measurement for Health-Care Applications", IEEE Digital Library Xplore, ISBN : 978-1-4673-2047-4, IEEE Catalog Number: CFP1221T-CDR, 2012.
4. M.Paranthaman, S.Palanivel Rajan, "Design of Implantable Antenna for Biomedical Applications", International Journal of Advanced Science and Technology, P-ISSN: 2005-4238, E-ISSN: 2207-6360, Vol. No.: 28, Issue No. 17, pp. 85-90, 2019.
5. Dr.S.Palanivel Rajan, Dr.C.Vivek, "Analysis and Design of Microstrip Patch Antenna for Radar Communication", Journal of Electrical Engineering & Technology, Online ISSN No.:

2093-7423, Print ISSN No.: 1975-0102, Vol. No.: 14, Issue : 2, DOI: 10.1007/s42835-018-00072-y, pp. 923–929, 2019.

6. Y. Makihara, A. Suzuki, D. Muramatsu, X. Li, and Y. Yagi, “Joint intensity and spatial metric learning for robust gait recognition,” in *Proc. 30th IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jul. 2017, pp. 5705–5715.
7. S.Palanivel Rajan, et.al., “Experimental Explorations on EOG Signal Processing for Real Time Applications in LabVIEW”, IEEE Digital Library Xplore, ISBN : 978-1-4673-2047-4, IEEE Catalog Number: CFP1221T-CDR, 2012.
8. Z. Wu, Y. Huang, L. Wang, X. Wang, and T. Tan, “A comprehensive study on cross-view gait based human identification with deep CNNs,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 39, no. 2, pp. 209–226, Mar. 2017.
9. Dr.S.Palanivel Rajan, Dr.C.Vivek, “Performance Analysis of Human Brain Stroke Detection System Using Ultra Wide Band Pentagon Antenna”, Sylwan Journal, ISSN No.: 0039-7660, Vol. No.: 164, Issue : 1, pp. 333–339, 2020.
10. S. Lombardi, K. Nishino, Y. Makihara, and Y. Yagi, “Two-point gait: Decoupling gait from body shape,” in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, Dec. 2013, pp. 1041–1048.
11. L. Niu, X. Xu, L. Chen, L. Duan, and D. Xu, “Action and event recognition in videos by learning from heterogeneous Web sources,” *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 28, no. 6, pp. 1290–1304, Jun. 2017.
12. Y. Zhou, X. Sun, Z.-J. Zha, and W. Zeng, “MICT: Mixed 3D/2D convolutional tube for human action recognition,” in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jun. 2018, pp. 449–458.
13. H. Liu, N. Shu, Q. Tang, and W. Zhang, “Computational model based
14. on neural network of visual cortex for human action recognition,” *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 29, no. 5, pp. 1427–1440, May 2018.
15. L. Sun, K. Jia, K. Chen, D.-Y. Yeung, B. E. Shi, and S. Savarese, “Lattice long short-term memory for human action recognition,” in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, Oct. 2017, pp. 2147–2156.
16. X. Chen, J. Weng, W. Lu, J. Xu, and J. Weng, “Deep manifold learning combined with convolutional neural networks for action recognition,” *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 29, no. 9, pp. 3938–3952, Sep. 2018.
17. J. Weng, M. Liu, X. Jiang, and J. Yuan, “Deformable pose traversal convolution for 3D action and gesture recognition,” in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, Sep. 2018, pp. 136–152.
18. J.-F. Hu, W.-S. Zheng, J. Pan, J. Lai, and J. Zhang, “Deep bilinear learning for RGB-D action recognition,” in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, Sep. 2018, pp. 335–351.
19. M. Guo, E. Chou, D.-A. Huang, S. Song, S. Yeung, and L. Fei-Fei, “Neural graph matching networks for fewshot 3D action recognition,” in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, Sep. 2018, pp. 653–669.
20. J. Wang, Z. Liu, Y. Wu, and J. Yuan, “Mining actionlet ensemble for action recognition with depth cameras,” in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jun. 2012, pp. 1290–1297.
21. Dr.S.Palanivel Rajan, L.Kavitha, “Automated retinal imaging system for detecting cardiac abnormalities using cup to disc ratio”, Indian Journal of Public Health Research & Development, Print ISSN: 0976-0245, Online ISSN: 0976-5506, Vol. No.: 10, Issue : 2, pp.1019-1024, DOI : 10.5958/0976-5506.2019.00430.3, 2019.

22. T.Abirami, S.Palanivel Rajan, "Cataloguing and Diagnosis of WBC'S in Microscopic Blood SMEAR", International Journal of Advanced Science and Technology, P-ISSN: 2005-4238, E-ISSN: 2207-6360, Vol. 28, Issue No. 17, pp. 69-76, 2019.
23. Rajan S. P, Paranthaman M. Novel Method for the Segregation of Heart Sounds from Lung Sounds to Extrapolate the Breathing Syndrome. Biosc.Biotech.Res.Comm. 2019;12(4).DOI: 10.21786/bbrc/12.4/1, 2019.
24. M. E. Hussein, M. Torki, M. A. Gowayyed, and M. El-Saban, "Human action recognition using a temporal hierarchy of covariance descriptors on 3D joint locations," in *Proc. IJCAI*, vol. 13, 2013, pp. 2466–2472.
25. W. Li, L. Wen, M.-C. Chang, S. N. Lim, and S. Lyu, "Adaptive RNN tree for large-scale human action recognition," in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, Oct. 2017, pp. 1444–1452.
26. Rajan, S., & Paranthaman, M. (2019). Characterization of compact and efficient patch antenna with single inset feeding technique for wireless applications. *Journal of Applied Research and Technology*, 17(4).
27. P. Zhang, C. Lan, J. Xing, W. Zeng, J. Xue, and N. Zheng, "View adaptive recurrent neural networks for high performance human action recognition from skeleton data," in *Proc. IEEE Int. Conf. Comput. Vis. (ICCV)*, Oct. 2017, pp. 2117–2126.
28. S. Yan, Y. Xiong, and D. Lin, "Spatial temporal graph convolutional networks for skeleton-based action recognition," 2018, *arXiv:1801.07455*. [Online]. Available: <https://arxiv.org/abs/1801.07455>
29. Dr.S.Palanivel Rajan, "Design of Microstrip Patch Antenna for Wireless Application using High Performance FR4 Substrate", *Advances and Applications in Mathematical Sciences*, ISSN No.: 0974-6803, Vol. No.: 18, Issue : 9, pp. 819-837, 2019.
30. Q. Ke, M. Bennamoun, S. An, F. Sohel, and F. Boussaid, "A new representation of skeleton sequences for 3D action recognition," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jul. 2017, pp. 4570–4579.