

## **Chapter 6 : MULTI-VIEW HUMAN ACTIVITY RECOGNITION BASED ON MULTIPLE FEATURES**

---

This chapter addresses the problem of silhouette based human activity recognition. Most of the previous work on silhouette based human activity recognition focuses on recognition from a single view and ignores the issue of view invariance. In this chapter, a system framework has been presented to recognize a view invariant human activity recognition approach that uses both contour-based pose features from silhouettes and uniform rotation local binary patterns for view invariant activity representation. The framework is composed of three consecutive modules: (i) detecting and locating people by background subtraction, (ii) combined scale invariant contour-based pose features from silhouettes and uniform rotation invariant local binary patterns (LBP) are extracted, and (iii) finally, classifying activities of people by Multiclass Support vector machine (SVM) classifier. The rotation invariant nature of uniform LBP provides view invariant recognition of multi view human activities. We have tested our approach successfully in the indoor and outdoor environment results on six multi-view datasets namely: our own view point dataset, VideoWeb Multi-view dataset [166], i3DPost multi-view dataset [164], WVU multi-view human action recognition dataset [167], MSR action recognition database [165] and i3DPost multi-view dataset [164] for multiple human. The experimental results show that the proposed method of multi-view human activity recognition is robust, flexible and efficient.

### **6.1. Introduction:**

Recognition of human activities from multiple views is a popular area of research in the field of computer vision. It is the basis of many applications in video surveillance and monitoring, human-computer interactions, model-based compressions, video retrieval in

various situations [24, 170]. Although a large amount of work has been performed on activity recognition in the last few years but still it is an open and challenging problem.

The various issues and challenges involved in automatic human activity recognition from video sequences are as follows: (i) clutter backgrounds, (ii) stationary or non-stationary camera, (iii) scale variation, (iv) starting and ending state variation, (v) individual variations in appearance and cloths of people, and (vi) changes in light and view-point. These situations make the human activity recognition a challenging task. Most of the work on activity recognition are view dependent and deal with recognition from one fixed view. To account these problems, many activity recognition systems and frameworks have been developed [24, 150-152, 155, 170-172].

Activity recognition methods available in the literature can broadly be categorized into two groups: sensor based activity recognition and vision based activity recognition. In sensor based activity recognition methods, some smart sensory device is used to capture various activity signals for activity recognition. Vision based activity recognition methods use the spatial or temporal structure of an activity in order to recognize it. A recent survey on vision-based action representation and recognition methods can be found in [22]. Machine learning based and Silhouette based methods [22] are popular vision based approaches for human activity recognition in videos. The machine learning based approaches for activity recognition generally solve the problem of activity recognition as a classification problem and classify an activity into one of known activity classes. Silhouette based methods are good options for activity recognition in video and can be easily used because of their simplicity and robustness. In silhouette based human activity methods [155], activities are recognized based on the key poses in the image. It is a local representation of activities. It decomposes the image into smaller interest regions and describes each region as a separate feature.

An immediate advantage of above mentioned approaches is that they neither rely on explicit body part labeling, nor on explicit human detection and localization. Bobick *et al.* [45], used motion templates for recognizing the activities in a specific environment of aerobic exercise. They used MEI (motion energy image) for obtaining foreground and MHI (motion history image) for obtaining motion information in a view-specific environment. It does not produce good activity recognition accuracy in an outdoor environment. Moreover, this technique is capable of only identifying one activity in the scene with one actor at a time. To deal the issue of Bobick *et al.* [45], Weinland *et al.* [58] proposed a method for multi-view human action recognition using Motion History Volumes (MHV) in 3D motion template. In this method, computing, aligning and comparing MHVs of different actions performed by different people in a variety of viewpoints. Weinland *et al.* [173] proposed another approach for multi-view human activity recognition which based on action taxonomy. In this method, an action taxonomy created using segmented sequence then apply a standard hierarchical clustering method to segment the sequence for human action recognition. Ahmad *et al.* [52] proposed a multi-view human action recognition system using combined local–global (CLG) optic flow based motion and shape feature for recognition of human actions in daily life. In this approach, actions are modeled by using a set of multidimensional HMMs for multiple views using the combined features of the training action videos. Iosifidis *et al.* [174] also presented a method for view-independent human action recognition. In this method, author used a new feature space (dyneme space) for multi-view movement representation and Fuzzy distances to represent the human body postures in the dyneme space. In this paper, view identification problem is solved by exploiting the circular shift invariance property of the Discrete Fourier Transform (DFT). Iosifidis *et al.* [41] proposed a method for multi-view human action recognition using neural network. In this paper, Fuzzy

distances from human body posture prototypes are used to produce a time invariant action representation and Multilayer perceptron are used for human action classification. This method is slow as it is based on neural network framework. To account this problem, Iosifidis *et al.* [175] presented a view-independent human action recognition method that exploits a 3D action representation. In this approach, binary human body images are temporally concatenated in order to produce action volumes (AVs) which represent the action. Multi-view action representation is obtained by exploiting the circular shift invariance property of the Discrete Fourier Transform coefficients. Chaaraoui *et al.* [176] proposed a method which is based on sequences of key poses. In this approach, learning of key poses is modeled using K-means and Dynamic Time Warping for action categorization. This approach suffer the problem of fixed view. To deal this issue, Iosifidis *et al.* [177] presented a method for view-independent human action recognition. In this method, fuzzy vector quantization is used for determine the human body poses. Sparsity-based Learning Machine (SbLM) has been used for view-independent action video classification. Sharma *et al.* [47] have proposed a human activity recognition method which used motion history images and object shape information for different human activities in a video. The major disadvantage of this method is that it is view dependent and deals with activity recognition from one fixed view.

To deal with the issues mentioned in [41, 45, 173-177], in this chapter, we have combined contour-based pose features from silhouettes and uniform rotation invariant local binary patterns (LBP) feature to model human activities. At first, change detection based approach is used for background subtraction. In the second step, combined contour-based pose features from silhouettes and uniform rotation invariant local binary patterns (LBP) are extracted. The contour-based pose features from silhouettes find the different key poses for human activities such as bending, standing, sitting etc. To solve view

depend (e.g. single or fixed view) problem, we have used uniform rotation invariant local binary patterns (LBP) feature. The uniform rotation invariant local binary patterns (LBP) feature provides view-independent analysis of human activities and it possess good discriminating ability, therefore they are better suited for distinguishing different activities. Finally, in a third step, this feature has been classified using Multiclass support vector machine (SVM) classifier.

The rest of the chapter is organized as follows: Section 6.2 briefly explains the proposed method. Section 6.3 presents experimental results and discussions. Finally, conclusions are drawn in Section 6.4.

## **6.2. Methods and Models**

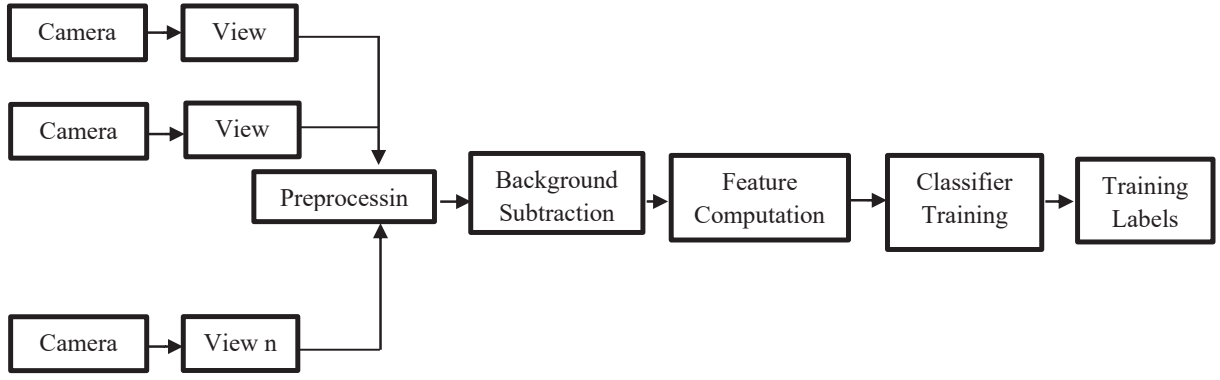
The proposed multi-view human activity recognition based on silhouette and uniform rotation invariant local binary patterns consists of three steps applied on given video frames which include: (i) detecting and locating people by background subtraction, (ii) combined scale invariant contour-based pose features from silhouettes and uniform rotation invariant local binary patterns (LBP) are extracted, and (iii) finally classifying activities of people by Multiclass Support vector machine (SVM) classifier. The proposed method consists of two phases testing and training as depicted in block diagram shown in Fig. 6.1 (a) and (b). In the training phase, we have given video data and apply above mention steps and finally training labels are given to the classifier. In the testing phase, testing video is performed with the help of training labels.

The algorithm for the proposed framework is as follows:

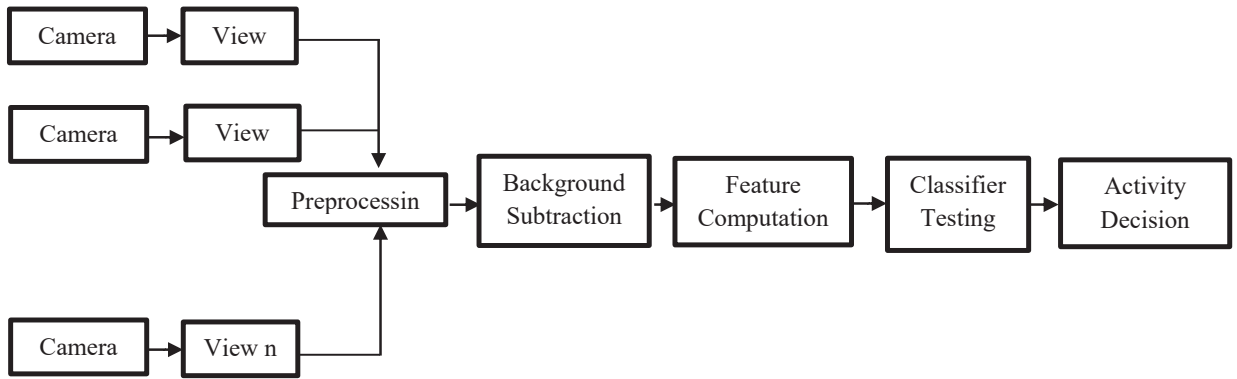
### **Algorithm**

**Step 1:** Capture video from the multiple cameras.

**Step 2:** Pre-processing of the capture video using background model creation (using Eq. 6.2 & 6.3).



(a) Training Phase



(b) Testing Phase

**Figure 6.1:** The block diagram of the proposed human activity recognition system

### Step 3: Feature Extraction

Computation of contour-based pose features from silhouettes for find out the different human key poses (using Eq. 6.12 & 6.13).

Computation of uniform rotation local binary patterns feature for view invariant recognition of multi-view human activities (using Eq. 6.17).

**Step 4:** Feature modeling and activity classification by using SVM classifier (using Eq. 6.20).

The various stages of the proposed method are discussed as follows:

### 6.2.1. Pre-processing

In the pre-processing steps, we extract foreground from the background. We then define the boundary from the foreground image sequence. Briefly, these are explained below:

#### 6.2.1.1. Background Subtraction

The proposed background subtraction method is based on the change detection and background modeling. The steps of background subtraction are as follows:

##### A. Frame difference

The difference between the current frame and the previous frame is calculated using change detection. Let  $f_n$  and  $f_{n-1}$  be the current frame and the previous frame at location (i, j). Instead of assigning a fixed *a priori* threshold  $V_{th,d}$  to each frame difference, this paper uses the fast Euler number computation technique [126] to automatically determine  $V_{th,d}$  from the video frame. The fast Euler numbers algorithm calculates the Euler number for every possible threshold with a single raster of the frame difference image using following equation:

$$E(i) = \frac{1}{4} [(q_1(i) - q_3(i) - 2q_d(i))] \quad (6.1)$$

where  $q_1$ ,  $q_3$ , and  $q_d$  is the quads (quad is a 2\*2 masks of bit cells) contained in the given image.

The output of the algorithm is an array of Euler numbers: one of each threshold value. The Zero Crossings find out the optimal threshold. Detailed algorithms for the fast Euler number computation method can be found in [126].

The frame differences  $WD_n(i, j)$  for respective frames are computed as:

*for every pixel location (i, j) in the co-ordinate of frame*

$$WD_n(i, j) = \begin{cases} 1 & \text{if } |f_n(i, j) - f_{n-1}(i, j)| > V_{th, WD} \\ 0 & \text{otherwise} \end{cases} \quad (6.2)$$

## B. Background model creation for segmentation

For background modeling, we have used frame difference, background registration, background difference, and background difference mask. The background modeling step is divided into five major phases. The first phase calculates the frame difference mask  $WD_n(i, j)$  which is obtained by difference between two consecutive frames as follows:

$$WD_n(i, j) = \begin{cases} 1 & \text{if } |Wf_n(i, j) - Wf_{n-1}(i, j)| \geq V_{th, WD} \\ 0, & \text{if } |Wf_n(i, j) - Wf_{n-1}(i, j)| < V_{th, WD} \end{cases} \quad (6.3)$$

$V_{th, WD}$  is a threshold determined automatically from the video frame by the fast Euler number computation method as explained in [126].

The second phase of dynamic background modeling maintains an up-to-date background buffer as well as background registration mask indicating whether the background information of a pixel is available or not. According to the frame difference mask of the past several frames, pixels that are not moving for a long time are considered as reliable background and registered in the background buffer. The background registration process uses the following two equations:

$$S_n(i, j) = \begin{cases} S_n(i, j) + 1 & \text{if } WD_n(i, j) = 0 \\ 0 & \text{if } WD_n(i, j) = 1 \end{cases} \quad (6.4)$$

$$\mu_n(i, j) = \begin{cases} f_n(i, j) & \text{if } S_n(i, j) \geq N_f \\ Undefined & \text{if } S_n(i, j) < N_f \end{cases} \quad (6.5)$$

where  $S_n(i, j)$  is a stationary index and  $\mu_n(i, j)$  is the background buffer value of a pixel with position  $(i, j)$  in the  $n^{\text{th}}$  frame. The initial values of  $S_n(i, j)$  and  $\mu_n(i, j)$  are

set to 0 and  $f_n(i, j)$ , respectively. If a pixel is masked as stationary for  $N_f$  successive frames (i.e., if the accumulated value in registration stationary index exceeds  $N_f$ ), then that pixel is classified as part of the background region. Here,  $N_f$  is set to 30 experimentally. According to our experiments,  $N_f$  may be set at a larger value for fast moving object.

In the third phase of background modeling, a registered background buffer pixel is updated using the following equation.

$$\text{if} \quad |f_n(i, j) - \mu_n(i, j)| < 2\sigma_n(i, j) \quad (6.6)$$

$$\text{then} \quad \begin{cases} \mu_n(i, j) = \chi\mu_{n-1}(i, j) + (1 - \chi)f_n(i, j) \\ \sigma_n^2(i, j) = \chi\sigma_{n-1}^2(i, j) + (1 - \chi)(f_n(i, j) - \mu_n(i, j))^2 \end{cases} \quad (6.7)$$

where  $\sigma_n(i, j)$  is the standard deviation of a pixel with position (i, j) in the  $n^{\text{th}}$  frame and  $\chi$  is the predefined constant and we considered four different sequences, and recorded 10 different observations over 800 frames for each of the sequences. This resulted in 50 samples of size 800 each. The test statistic was calculated for each of the samples and the value of  $\chi$  is set to 0.7

In the fourth phase of background modeling, we find the background difference mask with the help of background difference which distinguishes moving objects from the background, and its operation are shown as follows:

$$BD_{n,LL}(i, j) = |f_n(i, j) - \mu_n(i, j)| \quad (6.8)$$

$$BDM_n(i, j) = \begin{cases} 1, & BD_n(i, j) \geq V_{th,WD} \\ 0, & BD_n(i, j) < V_{th,WD} \end{cases} \quad (6.9)$$

where  $BD_{n,LL}(i, j)$  is the background difference and  $BDM_{n,LL}(i, j)$  is the background difference mask of a pixel with position  $(i, j)$  in the  $n^{\text{th}}$  frame. The threshold value  $V_{th,WD}$  is also automatically determined by the fast Euler number computation method [126].

In the fifth phase of background modeling, a background model is constructed using the frame difference, background registration, background difference, and background difference mask.

## 6.2.2. Multi-view Features Extraction

We use contour-based pose features and local binary pattern (LBP) feature for activities representation and classification. The foreground image sequence (which is obtained in section 6.2.1) is used to extract the contour-based pose features and local binary pattern (LBP).

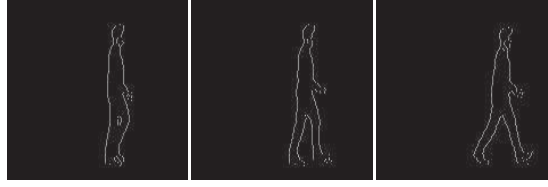
### 6.2.2.1. Contour-based Pose Feature

In this section, we find out the distance signal feature using contour points of the silhouette for different key poses (sitting, standing, sleeping etc.). We obtained a binary silhouette in section 6.2.1 by human silhouette extraction techniques, e.g. background subtraction.

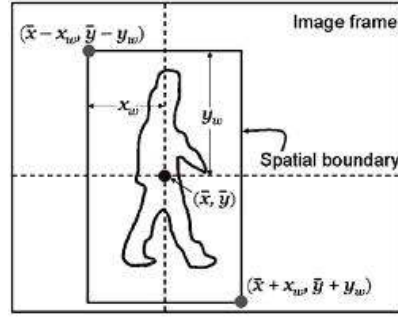
Let  $H$  be the binary silhouette image of an object. We determine its center of mass  $Cm = (\bar{x}, \bar{y})$  of the silhouette's contour points using [178] where

$$\bar{x} = \frac{\sum_{w=1}^n x_w}{n}, \quad \bar{y} = \frac{\sum_{w=1}^n y_w}{n} \quad (6.10)$$

and  $n$  is the number of silhouette pixels.



**Figure 6.2:** Sequence of key poses of walking activity in some selected frames



**Figure 6.3:** Activity boundary definitions

The distance signal  $D = \{d_1, d_2, d_3, \dots, d_n\}$  is generated by determining the Euclidean distance between each contour point and the centre of mass (see Fig 6.3.). Contour points should be considered always in the same order. For instance, the set of points can start at the most left point with equal y-axis value as the centre of mass, and follow a clockwise order.

$$d_i = \|C_m - a_i\|, \forall i \in [1, \dots, n]$$

(6.11)

In order to provide a uniform representation for varying image sizes and shapes,  $d$  is scaled to a constant size  $L$  such that

$$\hat{D}[i] = d \left\lceil i * \frac{n}{L} \right\rceil, \forall i \in [1, \dots, L] \quad (6.12)$$

where  $\lceil \cdot \rceil$  is the ceiling function.

Finally, the scaled distance vector  $\hat{D}$  is normalized to have unit sum:

$$\overline{D}[i] = \frac{\hat{D}[i]}{\sum_{i=1}^L \hat{D}[i]}, \forall i \in [1, \dots, L] \quad (6.13)$$

### 6.2.2.2. Uniform Rotation Invariant Local Binary Patterns (LBP)

In this section, we extract the Uniform rotation invariant local binary patterns (LBP) features from the background subtracted video which is obtained in section 6.2.1. Uniform rotation invariant local binary patterns (LBP) provide view invariant recognition of multi-view human activities. Using uniform patterns instead of all the possible patterns has produced better recognition results for human activity. The basic description of local binary patterns (LBP) is given below.

#### Local Binary Patterns (LBP)

A local binary pattern (LBP) feature can be constructed for a specific circular pixel neighborhood of radius  $R$ . The intensities of the  $P$  sample pixel points are compared in the circular neighborhood with the centre pixel in clockwise or anticlockwise direction (see Fig. 6.4).

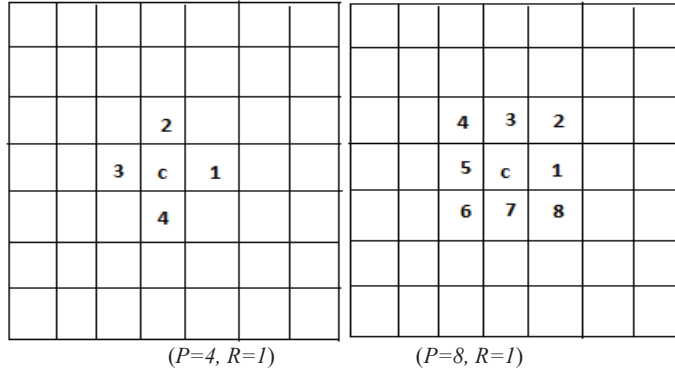


Figure 1 shows two 6x6 grids. The left grid is labeled  $(P=12, R=2)$  and the right grid is labeled  $(P=16, R=2)$ . Both grids show a sequence of numbers 1 through 11 in a specific pattern, with some cells containing 'c' or '16'.

|  |   |   |    |    |    |
|--|---|---|----|----|----|
|  |   |   |    |    |    |
|  |   | 5 | 4  | 3  |    |
|  | 6 |   |    |    | 2  |
|  | 7 |   | c  |    | 1  |
|  | 8 |   |    |    | 12 |
|  |   | 9 | 10 | 11 |    |
|  |   |   |    |    |    |

|  |    |    |    |    |    |
|--|----|----|----|----|----|
|  |    |    |    |    |    |
|  | 7  | 6  | 5  | 4  | 3  |
|  | 8  |    |    |    | 2  |
|  | 9  |    | c  |    | 1  |
|  | 10 |    |    |    | 16 |
|  | 11 | 12 | 13 | 14 | 15 |
|  |    |    |    |    |    |

After extracting LBP of each sample point in the image, value of each pixel in the image is replaced by a binary pattern. With the help of these considerations, the overall feature vector of the whole image, denoted by  $LBP_{P,R}$ , is given as below:

where  $(x, y)$  is the location of the centre pixel,  $g_c$  represent intensity of centre pixel,  $g_p$  represent intensity of neighborhood pixel and  $s(u)$  is defined as

Now, the feature vector  $LBP_{P,R}$  of the image is a histogram of the LBP of different pixels in the image. The starting size of the histogram is  $2^P$  because each possible LBP has been assigned a separate bin. Suppose, there are  $M$  regions in an image, then all histogram scan be merged into one histogram of size  $M \cdot 2^P$ .

Several modified versions of LBP [179] have been proposed for achieving rotation invariance and reducing the histogram dimension of the LBP. When the image is rotated, the gray value  $g_p$  will correspondingly move along the perimeter of the circle, so different

$LBP_{P,R}$  may be computed. To remove the effect of rotation, the modified version with rotation invariance is defined as follows

$$LBP_{P,R}^{ri}(x,y) = \min \{ ROR(LBP_{P,R}, i) \mid i = 0, 1, \dots, R-1 \} \quad (6.16)$$

Where  $ROR(LBP_{P,R}, i)$  performs a circular bit-wise right shift on the  $R$ -bit number  $LBP_{P,R}$  for  $i$  times.  $LBP_{P,R}^{ri}$  can have 36 different values when  $R=8$ , and the histogram dimension of  $LBP_{P,R}^{ri}$  over an image region is 36.

### Uniform Patterns

The uniform LBP are those LBP which have very few spatial transitions. Formally, uniform LBP have maximum two circular transitions between 0 and 1. For example, patterns 00000001 and 11111011 have only one and two transitions between 0 and 1 respectively, therefore they are uniform patterns.

### Uniform Local Binary Patterns (LBP) for Feature Extraction

The rotation invariant uniform local binary patterns (LBP) for feature extraction is defined as

$$LBP_{P,R}^{riu2} = \begin{cases} \sum_{p=0}^{P-1} s(g_p - g_c), & \text{if } U(LBP_{P,R}) \leq 2 \\ P+1, & \text{otherwise} \end{cases} \quad (6.17)$$

where  $U(LBP_{P,R}) = |s(g_{P-1} - g_c) - s(g_0 - g_c)| + \sum_{p=1}^{P-1} |s(g_p - g_c) - s(g_{p-1} - g_c)|$  is a rotation invariant operator with uniform patterns having at most two transitions between 0-1 bits.

In a circularly symmetric neighborhood of  $P$  pixels,  $P+1$  uniform pattern can be found.

Each pattern assigns a unique label to each pixel.

$$\text{and } s(u) = \begin{cases} 1, & u \geq 0 \\ 0, & u < 0 \end{cases} \quad (\text{given in Eq. 6.15})$$

$g_c$  = center pixel of background subtraction image (which is obtained in section 6.2.1)

and  $g_p$  = neighborhood pixel of background subtraction image (which is obtained in section 6.2.1).

### 6.2.3. Classifier Training and Testing

After having computed features from a video, the classifier is trained and tested with these video. To model and classify activities, we have used multi-class SVM classifiers. First training of classifier is performed. The obtained training labels are supplied into the classifier then after testing is performed. The activities in the testing video are performed with the help of training labels. Finally, different test labels have been obtained for the test video of human activities. Consider the pattern recognition problem of training samples  $(x_1, y_1), (x_2, y_2), \dots, (x_l, y_l)$ , where  $x_i, i = 1, 2, \dots, l$  is a vector and  $y_i \in \{1, 2, \dots, l\}$  represents the class of samples. The multi-class SVM [180] require the solution of the following optimization problem: minimize

$$\phi(\omega, \xi) = \frac{1}{2} \sum_{m=1}^k \omega_m \cdot \omega_m + C \sum_{i=1}^l \sum_{m \neq y_i} \xi_i^m \quad (6.18)$$

with constraints

$$(\omega_{y_i} \cdot x_i) + b_{y_i} \geq (\omega_m \cdot x_i) + b_m + 2 - \xi_i^m, \quad \xi_i^m \geq 0, \quad i = \{1, 2, \dots, l\}, m \in \{1, 2, \dots, k\} \setminus y_i \quad (6.19)$$

where  $C$  is the penalty parameter,  $l$  is the number of training data,  $k$  is the number of classes,  $y_i$  is the class of the  $i^{\text{th}}$  training data  $\omega$  points perpendicular to the separating hyper plane,  $b$  is the offset parameter to increase the margin, and  $\xi$  is the degree of misclassification of the datum  $x_i$ . This gives the decision function:

$$f(x) = \arg \max_{m=1, \dots, k} [\omega_m \cdot x + b_m] \quad (6.20)$$

### 6.3. Results, Analysis and Discussions

In this section, we perform the experiments and show results of the proposed method. We implemented human activity recognition method as described in Section 6.2 and tested on several datasets of human activity videos. Here, we present results for six representative publicly available human activity recognition video datasets – our own view point dataset, VideoWeb Multi-view dataset [166], i3DPost multi-view dataset [164], WVU multi-view human action recognition dataset [167], MSR action recognition database [165] and i3DPost multi-view dataset [164] for multiple human. Videos in these datasets have been captured at different rotation angle for multiple viewpoints. The experiments have been performed in Matlab R2013a window 7 environments on an Intel® Core™ i3 2.27 GHz machine with 4 GB RAM.

In our implementation, first we take the training videos and apply background subtraction according to the method described in Section 6.2.1. Secondly, scale invariant contour-based pose feature from silhouettes have been extracted. After that, extracting uniform rotation invariant local binary patterns (LBP) feature. Since it is rotation invariant, therefore provides robust results towards viewpoint changes. Uniform patterns provide good discriminating power. Lastly, multiclass SVM classifier has been employed for classification of activities in videos.

Six case studies of own view point dataset, VideoWeb Multi-view dataset [166], i3DPost multi-view dataset [164], WVU multi-view human action recognition dataset [167], MSR action recognition database [165] and i3DPost multi-view dataset [164] for multiple human are discussed here one by one. In all case studies, we have illustrated and tested the proposed method in comparison to Chaaraoui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], Iosifidis *et al.* [41], and Iosifidis *et al.* [175]. For quantitative

analysis of the proposed method and its comparative analysis with other methods correct recognition rate (CRR) is calculated which is defined as follows:

$$CRR = \frac{N_c}{N_a} \times 100 \text{ (in percentage)} \quad (6.21)$$

where  $N_c$  is the total number of correct recognition sequences while  $N_a$  is the number of total activity sequences.

### Experiment 1

In Fig. 6.5, we have shown results our own created database from different viewpoints. This database contains video of static human activities namely standing and two dynamic activities namely bending and jogging in different view direction. These videos are taken in real outdoor environment. From the observation of this figure, it is clear that the proposed method is well capable of recognizing these static and dynamic activities. Moreover, there is some little movement in each activity, i.e. pose of human object does not remain still for all the time. Direction of each human object also changes in different frames. Therefore, the proposed method is pose invariant and frontal view is not necessary for recognition of objects and suits for recognition of objects with frontal as well as side view. The proposed method is capable of recognizing the activity at these different viewing angles correctly. The proposed method is also robust towards different rotations of the activity.

Now, we show quantitative results of the proposed method and compare them with other existing methods in terms of confusion matrix. The other methods are Chaaraoui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], Iosifidis *et al.* [41], and Iosifidis *et al.* [175].



(a) Bending



(b) Jogging



(c) Standing

**Figure 6.5:** Recognition of activities in our own database (a) Bending (b) Jogging (c) standing in different views.

The confusion matrix of different activity for different methods has been shown in Table 6.1. After observing these tables, we see that the diagonal values are the highest for the proposed method in each case.

**Table 6.1:** Confusion matrices for the proposed and other methods

| <b>Recognized Instances</b><br>→     | <b>Bending</b> | <b>Jogging</b> | <b>Walking</b> |
|--------------------------------------|----------------|----------------|----------------|
| <b>Total Instances</b><br>↓          |                |                |                |
| <b>Proposed method</b>               |                |                |                |
| <b>Bending</b>                       | 0.98           | 0.02           | 0              |
| <b>Jogging</b>                       | 0              | .99            | 0.01           |
| <b>Walking</b>                       | 0              | 0              | 1              |
| <b>Chaaraoui <i>et al.</i> [176]</b> |                |                |                |
| <b>Bending</b>                       | 0.85           | 0.03           | 0.12           |
| <b>Jogging</b>                       | 0.15           | 0.80           | 0.05           |
| <b>Walking</b>                       | 0.02           | 0.10           | 0.82           |
| <b>Bobick <i>et al.</i> [45]</b>     |                |                |                |
| <b>Bending</b>                       | 0.75           | 0.10           | 0.15           |
| <b>Jogging</b>                       | 0.15           | 0.70           | 0.15           |
| <b>Walking</b>                       | 0.12           | 0.11           | 0.77           |
| <b>Ahmad <i>et al.</i> [52]</b>      |                |                |                |
| <b>Bending</b>                       | 0.79           | 0.20           | 0.01           |
| <b>Jogging</b>                       | 0.10           | 0.82           | 0.08           |
| <b>Walking</b>                       | 0.12           | 0              | 0.88           |
| <b>Iosifidis <i>et al.</i> [177]</b> |                |                |                |
| <b>Bending</b>                       | 0.94           | 0.06           | 0              |
| <b>Jogging</b>                       | 0.03           | 0.95           | 0.02           |
| <b>Walking</b>                       | 0.02           | 0              | 0.98           |
| <b>Weinland <i>et al.</i> [58]</b>   |                |                |                |
| <b>Bending</b>                       | 0.82           | 0.18           | 0              |
| <b>Jogging</b>                       | 0.10           | 0.84           | 0.06           |
| <b>Walking</b>                       | 0.14           | 0              | 0.86           |
| <b>Iosifidis <i>et al.</i> [174]</b> |                |                |                |
| <b>Bending</b>                       | 0.90           | 0.05           | 0.05           |
| <b>Jogging</b>                       | 0.06           | 0.94           | 0              |
| <b>Walking</b>                       | 0.09           | 0              | 0.91           |
| <b>Weinland <i>et al.</i> [173]</b>  |                |                |                |
| <b>Bending</b>                       | 0.78           | 0.12           | 0.10           |
| <b>Jogging</b>                       | 0.20           | 0.80           | 0              |
| <b>Walking</b>                       | 0.22           | 0              | 0.78           |
| <b>Iosifidis <i>et al.</i> [41]</b>  |                |                |                |
| <b>Bending</b>                       | 0.96           | 0.04           | 0              |
| <b>Jogging</b>                       | 0              | 0.96           | 0.04           |
| <b>Walking</b>                       | 0.02           | 0.02           | 0.94           |
| <b>Iosifidis <i>et al.</i> [175]</b> |                |                |                |
| <b>Bending</b>                       | 0.98           | 0.02           | 0              |
| <b>Jogging</b>                       | 0.02           | 0.94           | 0.02           |
| <b>Walking</b>                       | 0              | 0.05           | 0.95           |

**Table 6.2:** Recognition results over our own activity recognition dataset

| Method                        | Accuracy (%) |
|-------------------------------|--------------|
| Proposed Method               | 99           |
| Charaoui <i>et al.</i> [176]  | 82           |
| Bobick <i>et al.</i> [45]     | 74           |
| Ahmad <i>et al.</i> [52]      | 83           |
| Iosifidis <i>et al.</i> [177] | 95.66        |
| Weinland <i>et al.</i> [58]   | 84           |
| Iosifidis <i>et al.</i> [174] | 91.66        |
| Weinland <i>et al.</i> [173]  | 78.66        |
| Iosifidis <i>et al.</i> [41]  | 95.33        |
| Iosifidis <i>et al.</i> [175] | 95.66        |

A comparison of recognition accuracy of different methods has been shown in Table 6.2 (calculate by using Eq. 6.21). Higher the value, higher will be the recognition accuracy.

From these confusion matrices and recognition results, it can be observed that the performance of the proposed method is better in comparison to other existing methods.

The recognition accuracy of the proposed method is better than other methods.

## Experiment 2

We have demonstrated results of the proposed method for VideoWeb Multi-view dataset [166]. VideoWeb dataset involves up to 10 actors interacting in various ways (with each other, with vehicles or with facilities). The activities are: waving, boxing, clapping, jogging, running and walking. It consists of about 2.5 hours of video recorded from a minimum of 4 and a maximum of 8 cameras. Each video is recorded by a camera network whose number of cameras depends on the type of scene.



(a) Boxing



(b) Clapping



(c) Jogging



(d) Running



(e) Walking

**Figure 6.6:** Recognition of Activities in VideoWeb Multi-view dataset [166] (a) Boxing (b) Clapping (c) Jogging (d) Running (e) Walking

From Fig. 6.6, it can be observed that the person is performing different activity such as boxing, clapping, jogging, running, and walking at different viewing angles. From Fig. 6.6, it is concluded that the pose of human object does not remain still for all the time. Direction of each human object also changes in different frames. Therefore, the proposed method is pose invariant and frontal view is not necessary for recognition of objects and suits for recognition of objects with frontal as well as side view. These visual results show that the obtained results are accurate and the proposed method provide proper recognition results for VideoWeb Multi-view dataset [166].

Now, we show quantitative results of the proposed method and compare them with other existing methods in terms of confusion matrix. The other methods are Chaaaraoui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], Iosifidis *et al.* [41], and Iosifidis *et al.* [175].

**Table 6.3:** Confusion matrices for the proposed and other methods

| Recognized Instances<br>→<br>Total Instances↓ | Boxing | Clapping | Jogging | Running | Walking |
|-----------------------------------------------|--------|----------|---------|---------|---------|
| <b>Proposed method</b>                        |        |          |         |         |         |
| Boxing                                        | 1      | 0        | 0       | 0       | 0       |
| Clapping                                      | 0      | 0.97     | 0.02    | 0       | 0.01    |
| Jogging                                       | 0      | 0.02     | 0.98    | 0       | 0       |
| Running                                       | 0      | 0        | 0       | 1       | 00      |
| Walking                                       | 0      | 0        | 0       | 0       | 1       |
| <b>Bobick <i>et al.</i> [45]</b>              |        |          |         |         |         |
| Boxing                                        | 0.72   | 0.10     | 0       | 0.18    | 0       |
| Clapping                                      | 0      | 0.74     | 0.20    | 0       | 0.06    |
| Jogging                                       | 0.15   | 0        | 0.70    | 0.12    | 0.03    |
| Running                                       | 0.10   | 0.10     | 0.09    | 0.71    | 0       |
| Walking                                       | 0.15   | 0        | 0       | 0.15    | 0.70    |
| <b>Chaaroui <i>et al.</i> [176]</b>           |        |          |         |         |         |
| Boxing                                        | 0.78   | 0        | 0.15    | 0       | 0.07    |
| Clapping                                      | 0.10   | 0.80     | 0       | 0.10    | 0       |
| Jogging                                       | 0.12   | 0.07     | 0.81    | 0       | 0       |
| Running                                       | 0.13   | 0        | 0.11    | 0.76    | 0       |
| Walking                                       | 0      | 0.15     | 0       | 0.05    | 0.80    |
| <b>Ahmad <i>et al.</i> [52]</b>               |        |          |         |         |         |
| Boxing                                        | 0.86   | 0.10     | 0       | 0       | 0.04    |
| Clapping                                      | 0.03   | 0.84     | 0.12    | 0.01    | 0       |
| Jogging                                       | 0.15   | 0        | 0.83    | 0       | 0.02    |
| Running                                       | 0.10   | 0.10     | 0       | 0.80    | 0       |
| Walking                                       | 0      | 0        | 0.15    | 0       | 0.85    |
| <b>Iosifidis <i>et al.</i> [177]</b>          |        |          |         |         |         |
| Boxing                                        | 0.98   | 0.02     | 0       | 0       | 0       |
| Clapping                                      | 0.05   | 0.95     | 0       | 0       | 0       |
| Jogging                                       | 0      | 0        | 0.95    | 0.05    | 0       |
| Running                                       | 0.01   | 0        | 0       | 0.93    | 0.06    |
| Walking                                       | 0      | 0.04     | 0       | 0       | 0.96    |
| <b>Weinland <i>et al.</i> [58]</b>            |        |          |         |         |         |
| Boxing                                        | 0.88   | 0.12     | 0       | 0       | 0       |
| Clapping                                      | 0      | 0.82     | 0.10    | 0.08    | 0       |
| Jogging                                       | 0.13   | 0        | 0.87    | 0       | 0       |
| Running                                       | 0      | 0        | 0       | 0.88    | 0.12    |
| Walking                                       | 0.12   | 0.04     | 0       | 0       | 0.84    |
| <b>Iosifidis <i>et al.</i> [174]</b>          |        |          |         |         |         |
| Boxing                                        | 0.91   | 0.09     | 0       | 0       | 0       |
| Clapping                                      | 0.06   | 0.94     | 0       | 0       | 0       |
| Jogging                                       | 0      | 0        | 0.91    | 0.09    | 0       |
| Running                                       | 0      | 0        | 0.05    | 0.90    | 0.05    |
| Walking                                       | 0      | 0        | 0.09    | 0       | 0.91    |

| <b>Weinland <i>et al.</i> [173]</b>  |      |      |      |      |      |
|--------------------------------------|------|------|------|------|------|
| <b>Boxing</b>                        | 0.77 | 0.20 | 0.03 | 0    | 0    |
| <b>Clapping</b>                      | 0.20 | 0.80 | 0    | 0    | 0    |
| <b>Jogging</b>                       | 0    | 0    | 0.82 | 0.18 | 0    |
| <b>Running</b>                       | 0.18 | 0    | 0    | 0.82 | 0    |
| <b>Walking</b>                       | 0.12 | 0    | 0.10 | 0    | 0.78 |
| <b>Iosifidis <i>et al.</i> [41]</b>  |      |      |      |      |      |
| <b>Boxing</b>                        | 0.94 | 0.06 | 0    | 0    | 0    |
| <b>Clapping</b>                      | 0.03 | 0.97 | 0    | 0    | 0    |
| <b>Jogging</b>                       | 0    | 0    | 0.90 | 0.05 | 0.05 |
| <b>Running</b>                       | 0.09 | 0    | 0    | 0.91 | 0    |
| <b>Walking</b>                       | 0    | 0    | 0    | 0.10 | 0.90 |
| <b>Iosifidis <i>et al.</i> [175]</b> |      |      |      |      |      |
| <b>Boxing</b>                        | 0.97 | 0.03 | 0    | 0    | 0    |
| <b>Clapping</b>                      | 0    | 0.96 | 0    | 0.04 | 0    |
| <b>Jogging</b>                       | 0    | 0    | 0.98 | 0    | 0.02 |
| <b>Running</b>                       | 0    | 0    | 0.03 | 0.97 | 0    |
| <b>Walking</b>                       | 0.07 | 0    | 0    | 0    | 0.93 |

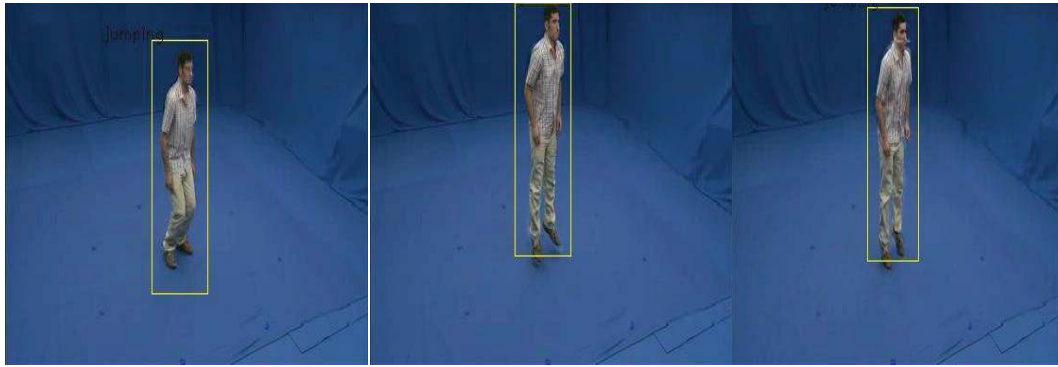
**Table 6.4:** Recognition results over VideoWeb action recognition dataset [166]

| <b>Method</b>                 | <b>Accuracy (%)</b> |
|-------------------------------|---------------------|
| Proposed Method               | 99                  |
| Charaoui <i>et al.</i> [176]  | 64.4                |
| Bobick <i>et al.</i> [45]     | 71.4                |
| Ahmad <i>et al.</i> [52]      | 83.6                |
| Iosifidis <i>et al.</i> [177] | 95.4                |
| Weinland <i>et al.</i> [58]   | 85.8                |
| Iosifidis <i>et al.</i> [174] | 91.4                |
| Weinland <i>et al.</i> [173]  | 79.8                |
| Iosifidis <i>et al.</i> [41]  | 92.4                |
| Iosifidis <i>et al.</i> [175] | 96.2                |

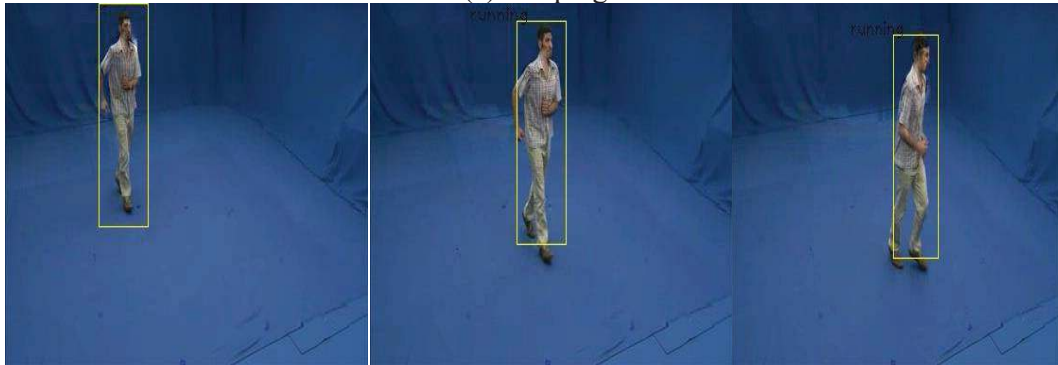
The confusion matrix of different activities for different methods has been shown in Table 6.3. After observing these tables, we see that the diagonal values are the highest for the proposed method in each case. A comparison of recognition accuracy of different methods has been shown in Table 6.4. From these confusion matrices and recognition results, it can be observed that the performance and recognition accuracy of the proposed method is better in comparison to other existing methods.

### Experiment 3

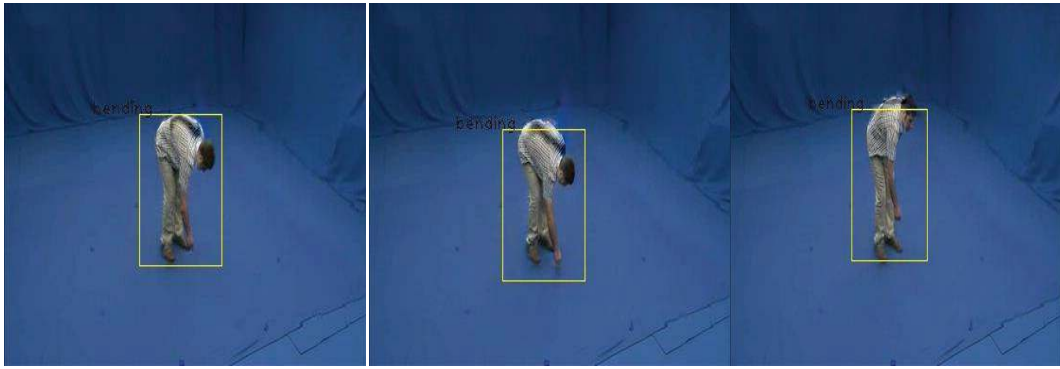
We have shown activity recognition with standard i3DPost dataset which is a multi-view dataset [164]. In this dataset, 8 people performing 13 actions (walking, running, jumping, bending, hand-waving, jumping in place, sitting-stand up, running-falling, walking-sitting, running-jumping-walking, handshaking, pulling, and facial-expressions) each one. In Fig.6.7, six different activities have been performed on multi-view. These activities have been performed with the help of 5 cameras placed at different viewing angles and activities have been captured simultaneously with these cameras. These visual results show that the obtained results are accurate and the proposed method provide proper recognition results for this set of videos also.



(a) Jumping



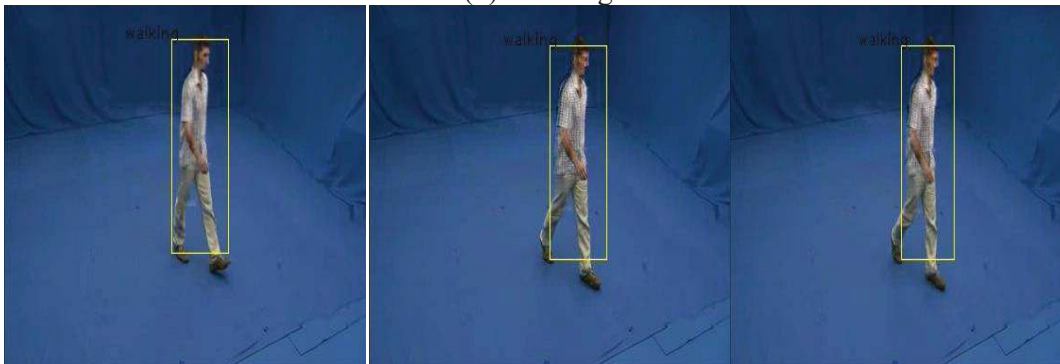
(b) Running



(c) Bending



(d) Standing



(e) Walking



(f) Sitting

**Figure 6.7:** Recognition of Activities in i3DPost multi-view dataset [164] (a) Jumping (b) Running (c) Bending (d) Standing (e) Walking (f) Sitting

Now, quantitative results have been shown for i3DPost multi-view dataset [164] in Tables 6.5-6.6. These confusion matrices and recognition results in Tables 6.5 & 6.6 indicate that the proposed method performs better than other methods.

**Table 6.5:** Confusion matrix for the proposed method and other methods

| <b>Recognized<br/>Instances</b><br><b>→</b><br><b>Total<br/>Instances</b><br><b>↓</b> | <b>Jumping</b> | <b>Running</b> | <b>Bending</b> | <b>Standing</b> | <b>Walking</b> | <b>Sitting</b> |
|---------------------------------------------------------------------------------------|----------------|----------------|----------------|-----------------|----------------|----------------|
| <b>Proposed method</b>                                                                |                |                |                |                 |                |                |
| <b>Jumping</b>                                                                        | 0.98           | 0.01           | 0              | 0.01            | 0              | 0              |
| <b>Running</b>                                                                        | 0              | 1              | 0              | 0               | 0              | 0              |
| <b>Bending</b>                                                                        | 0              | 0              | 0.99           | 0               | 0.01           | 0              |
| <b>Standing</b>                                                                       | 0.02           | 0              | 0              | 0.98            | 0              | 0              |
| <b>Walking</b>                                                                        | 0              | 0              | 0              | 0               | 1              | 0              |
| <b>Sitting</b>                                                                        | 0              | 0.01           | 0              | 0               | 0.01           | 0.98           |
| <b>Chaaroui et al. [176]</b>                                                          |                |                |                |                 |                |                |
| <b>Jumping</b>                                                                        | 0.80           | 0.10           | 0.05           | 0               | 0.05           | 0              |
| <b>Running</b>                                                                        | 0.10           | 0.85           | 0.05           | 0               | 0              | 0              |
| <b>Bending</b>                                                                        | 0.08           | 0.05           | 0.82           | 0               | 0.05           | 0              |
| <b>Standing</b>                                                                       | 0              | 0              | 0.10           | 0.84            | 0              | 0.06           |
| <b>Walking</b>                                                                        | 0              | 0              | 0              | 0.10            | 0.80           | 0.10           |
| <b>Sitting</b>                                                                        | 0              | 0.10           | 0              | 0.01            | 0.08           | 0.81           |
| <b>Bobick et al. [45]</b>                                                             |                |                |                |                 |                |                |
| <b>Jumping</b>                                                                        | 0.75           | 0.10           | 0.03           | 0.10            | 0              | 0.02           |
| <b>Running</b>                                                                        | 0.07           | 0.78           | 0.12           | 0               | 0.03           | 0              |
| <b>Bending</b>                                                                        | 0              | 0.10           | 0.85           | 0               | 0.05           | 0              |
| <b>Standing</b>                                                                       | 0              | 0.10           | 0.07           | 0.80            | 0.01           | 0.02           |
| <b>Walking</b>                                                                        | 0.04           | 0.02           | 0.07           | 0.07            | 0.80           | 0              |
| <b>Sitting</b>                                                                        | 0.12           | 0              | 0              | 0.04            | 0              | 0.84           |
| <b>Ahmad et al. [52]</b>                                                              |                |                |                |                 |                |                |
| <b>Jumping</b>                                                                        | 0.87           | 0.05           | 0.05           | 0.03            | 0              | 0              |
| <b>Running</b>                                                                        | 0.06           | 0.90           | 0              | 0.04            | 0              | 0              |
| <b>Bending</b>                                                                        | 0.10           | 0              | 0.86           | 0               | 0.02           | 0.02           |
| <b>Standing</b>                                                                       | 0.05           | 0.07           | 0              | 0.84            | 0.02           | 0.02           |
| <b>Walking</b>                                                                        | 0              | 0.05           | 0.05           | 0.01            | 0.88           | 0.01           |
| <b>Sitting</b>                                                                        | 0              | 0.05           | 0              | 0.06            | 0.02           | 0.87           |
| <b>Iosifidis et al. [177]</b>                                                         |                |                |                |                 |                |                |
| <b>Jumping</b>                                                                        | 0.98           | 0.02           | 0              | 0               | 0              | 0              |
| <b>Running</b>                                                                        | 0              | 0.97           | 0              | 0.03            | 0              | 0              |
| <b>Bending</b>                                                                        | 0              | 0              | 0.99           | 0               | 0.01           | 0              |
| <b>Standing</b>                                                                       | 0              | 0              | 0              | 0.95            | 0              | 0.05           |
| <b>Walking</b>                                                                        | 0.02           | 0              | 0              | 0               | 0.98           | 0              |
| <b>Sitting</b>                                                                        | 0              | 0              | 0.04           | 0               | 0              | 0.96           |
| <b>Weinland et al. [58]</b>                                                           |                |                |                |                 |                |                |

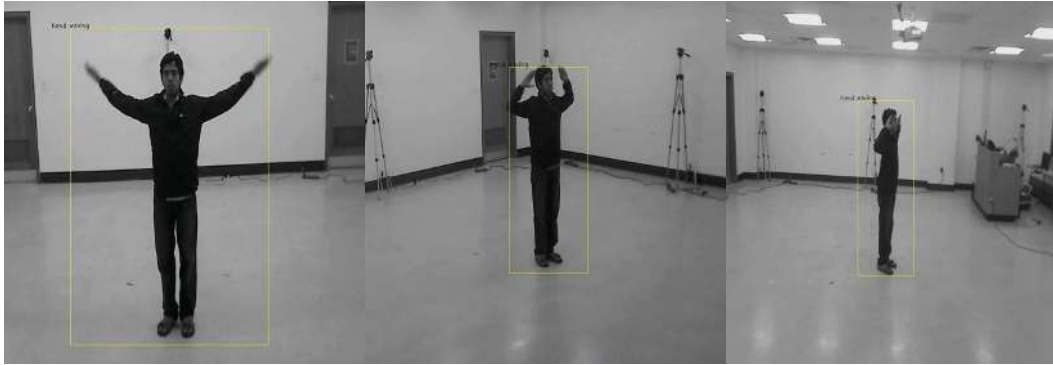
|                               |      |      |      |      |      |      |
|-------------------------------|------|------|------|------|------|------|
| <b>Jumping</b>                | 0.80 | 0.10 | 0    | 0    | 0    | 0    |
| <b>Running</b>                | 0    | 0.86 | 0.04 | 0    | 0    | 0.10 |
| <b>Bending</b>                | 0    | 0    | 0.88 | 0    | 0.12 | 0    |
| <b>Standing</b>               | 0    | 0    | 0.17 | 0.83 | 0    | 0    |
| <b>Walking</b>                | 0    | 0    | 0    | 0.18 | 0.82 | 0    |
| <b>Sitting</b>                | 0.16 | 0    | 0    | 0    | 0    | 0.84 |
| <b>Iosifidis et al. [174]</b> |      |      |      |      |      |      |
| <b>Jumping</b>                | 0.90 | 0    | 0    | 0.10 | 0    | 0    |
| <b>Running</b>                | 0.06 | 0.94 | 0    | 0    | 0    | 0    |
| <b>Bending</b>                | 0    | 0    | 0.93 | 0    | 0.07 | 0    |
| <b>Standing</b>               | 0    | 0    | 0    | 0.92 | 0    | 0.08 |
| <b>Walking</b>                | 0.05 | 0    | 0    | 0    | 0.95 | 0    |
| <b>Sitting</b>                | 0    | 0    | 0    | 0.09 | 0    | 0.91 |
| <b>Weinland et al. [173]</b>  |      |      |      |      |      |      |
| <b>Jumping</b>                | 0.76 | 0.12 | 0.12 | 0    | 0    | 0    |
| <b>Running</b>                | 0    | 0.80 | 0    | 0    | 0    | 0.20 |
| <b>Bending</b>                | 0    | 0    | 0.82 | 0.18 | 0    | 0    |
| <b>Standing</b>               | 0    | 0    | 0.20 | 0.79 | 0.01 | 0    |
| <b>Walking</b>                | 0    | 0    | 0    | 0    | 0.80 | 0.20 |
| <b>Sitting</b>                | 0    | 0    | 0    | 0    | 0.19 | 0.81 |
| <b>Iosifidis et al. [41]</b>  |      |      |      |      |      |      |
| <b>Jumping</b>                | 0.96 | 0    | 0    | 0.04 | 0    | 0    |
| <b>Running</b>                | 0    | 0.95 | 0    | 0    | 0.05 | 0    |
| <b>Bending</b>                | 0.06 | 0    | 0.94 | 0    | 0    | 0    |
| <b>Standing</b>               | 0    | 0    | 0.04 | 0.96 | 0    | 0    |
| <b>Walking</b>                | 0    | 0    | 0    | 0    | 0.95 | 0.05 |
| <b>Sitting</b>                | 0    | 0.04 | 0    | 0    | 0    | 0.96 |
| <b>Iosifidis et al. [175]</b> |      |      |      |      |      |      |
| <b>Jumping</b>                | 0.98 | 0.02 | 0    | 0    | 0    | 0    |
| <b>Running</b>                | 0    | 0.96 | 0.04 | 0    | 0    | 0    |
| <b>Bending</b>                | 0    | 0    | 0.99 | 0    | 0.01 | 0    |
| <b>Standing</b>               | 0    | 0    | 0    | 0.95 | 0.05 | 0    |
| <b>Walking</b>                | 0.06 | 0    | 0    | 0    | 0.94 | 0    |
| <b>Sitting</b>                | 0    | 0    | 0    | 0.08 | 0    | 0.92 |

**Table 6.6:** Recognition results over the i3DPost multi-view dataset [164]

| <b>Method</b>          | <b>Accuracy (%)</b> |
|------------------------|---------------------|
| Proposed Method        | 98.83               |
| Charaoui et al. [176]  | 82                  |
| Bobick et al. [45]     | 80.33               |
| Ahmad et al. [52]      | 87                  |
| Iosifidis et al. [177] | 97.16               |
| Weinland et al. [58]   | 83.83               |
| Iosifidis et al. [174] | 92.5                |
| Weinland et al. [173]  | 79.66               |
| Iosifidis et al. [41]  | 95.33               |
| Iosifidis et al. [175] | 95.66               |

## Experiment 4

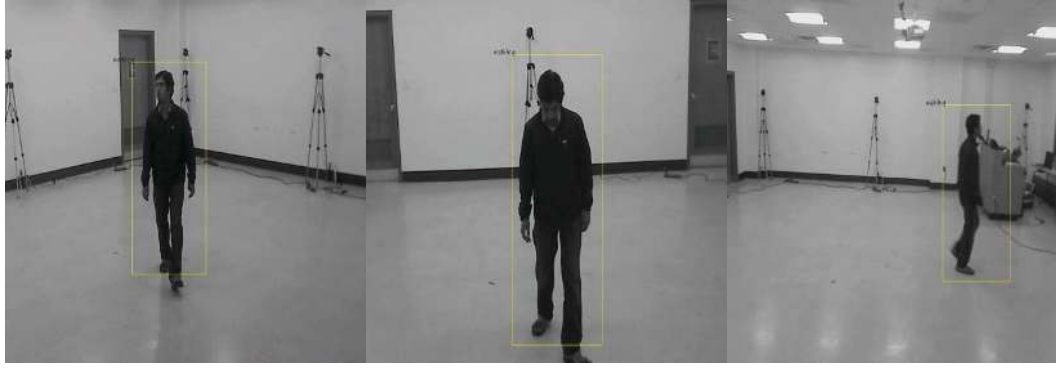
In Fig.6.8, we have shown activity recognition with WVU multi-view human action recognition dataset [167]. This database includes different activities hand waving, clapping, jumping, jogging, bowling, throwing, pickup, and kicking. WVU multi-view human action recognition dataset [167] has been sorted based on the 8 views. For each view, action sequences performed by different subjects are provided. In Fig.6.8, it is easily concluded that the proposed method is invariant with respect to pose of the human object and also a frontal view is not necessary for recognition of objects and gives satisfactory results for human objects with frontal as well as side view.



(a) Hand Waving



(b) Hand Clapping



(c) Walking

**Figure 6.8:** Recognition of Activities in WVU multi-view human action recognition dataset [167] (a) Hand waving (b) Hand Clapping (c) Walking

Now, quantitative results have been shown WVU multi-view human action recognition dataset [167] in Tables 6.7-6.8.

These confusion matrices and recognition results presented in Tables 6.7 and 6.8 show that the accuracy of the proposed method is better than the other existing methods. Each confusion matrix shows the performance of a particular method for the chosen dataset.

**Table 6.7:** Confusion matrix for the proposed method and other methods

| Recognized Instances<br>→<br>Total Instances↓ | Hand Waving | Hand-clapping | Walking |
|-----------------------------------------------|-------------|---------------|---------|
| <b>Proposed method</b>                        |             |               |         |
| Hand Waving                                   | 1           | 0             | 0       |
| Hand-clapping                                 | 0           | 1             | 0       |
| Walking                                       | 0           | 0.02          | 0.98    |
| <b>Chaaroui et al. [176]</b>                  |             |               |         |
| Hand Waving                                   | 0.72        | 0.28          | 0       |
| Hand-clapping                                 | 0.30        | 0.70          | 0       |
| Walking                                       | 0.30        | 0.01          | 0.69    |
| <b>Bobick et al. [45]</b>                     |             |               |         |
| Hand Waving                                   | 0.74        | 0.26          | 0       |
| Hand-clapping                                 | 0.24        | 0.76          | 0       |
| Walking                                       | 0           | 0.23          | 0.77    |
| <b>Ahmad et al. [52]</b>                      |             |               |         |
| Hand Waving                                   | 0.84        | 0.16          | 0       |
| Hand-clapping                                 | 0           | 0.81          | 0.19    |
| Walking                                       | 0.07        | 0.10          | 0.83    |
| <b>Iosifidis et al. [177]</b>                 |             |               |         |
| Hand Waving                                   | 0.97        | 0             | 0.03    |
| Hand-clapping                                 | 0.07        | 0.93          | 0       |
| Walking                                       | 0           | 0.08          | 0.92    |
| <b>Weinland et al. [58]</b>                   |             |               |         |
| Hand Waving                                   | 0.88        | 0.12          | 0       |
| Hand-clapping                                 | 0           | 0.85          | 0.15    |
| Walking                                       | 0           | 0.13          | 0.87    |
| <b>Iosifidis et al. [174]</b>                 |             |               |         |
| Hand Waving                                   | 0.93        | 0.07          | 0       |
| Hand-clapping                                 | 0.09        | 0.91          | 0       |
| Walking                                       | 0           | 0.07          | 0.93    |
| <b>Weinland et al. [173]</b>                  |             |               |         |
| Hand Waving                                   | 0.78        | 0             | 0.22    |
| Hand-clapping                                 | 0.09        | 0.81          | 0.10    |
| Walking                                       | 0           | 0.18          | 0.82    |
| <b>Iosifidis et al. [41]</b>                  |             |               |         |
| Hand Waving                                   | 0.94        | 0.06          | 0       |
| Hand-clapping                                 | 0.05        | 0.95          | 0       |
| Walking                                       | 0           | 0.08          | 0.92    |
| <b>Iosifidis et al. [175]</b>                 |             |               |         |
| Hand Waving                                   | 0.97        | 0.03          | 0       |
| Hand-clapping                                 | 0.05        | 0.95          | 0       |
| Walking                                       | 0.05        | 0             | 0.95    |

**Table 6.8:** Recognition results over the WVU action recognition dataset [167]

| Method                         | Accuracy (%) |
|--------------------------------|--------------|
| Proposed Method                | 99.33        |
| Chaaaraoui <i>et al.</i> [176] | 70.33        |
| Bobick <i>et al.</i> [45]      | 75.66        |
| Ahmad <i>et al.</i> [52]       | 85           |
| Iosifidis <i>et al.</i> [177]  | 94           |
| Weinland <i>et al.</i> [58]    | 86.66        |
| Iosifidis <i>et al.</i> [174]  | 92.33        |
| Weinland <i>et al.</i> [173]   | 80.33        |
| Iosifidis <i>et al.</i> [41]   | 93.66        |
| Iosifidis <i>et al.</i> [175]  | 95.66        |

### Experiment 5

In this section, we demonstrate results of the proposed method for MSR action recognition database [165]. MSR Action dataset contains 16 video sequences and total 63 actions: 14 hand clapping, 24 hand-waving and 25 boxing, performed by 10 subjects. Each sequence contains multiple types of actions. Some sequences contain actions performed by different people. There are both indoor and outdoor scenes. All of the video sequences are captured with clutter and moving backgrounds. Each video is of low resolution 320 x 240 and frame rate 15 frames per second. Their lengths are between 32 to 76 seconds. Qualitative recognition results are shown in Fig. 6.9.

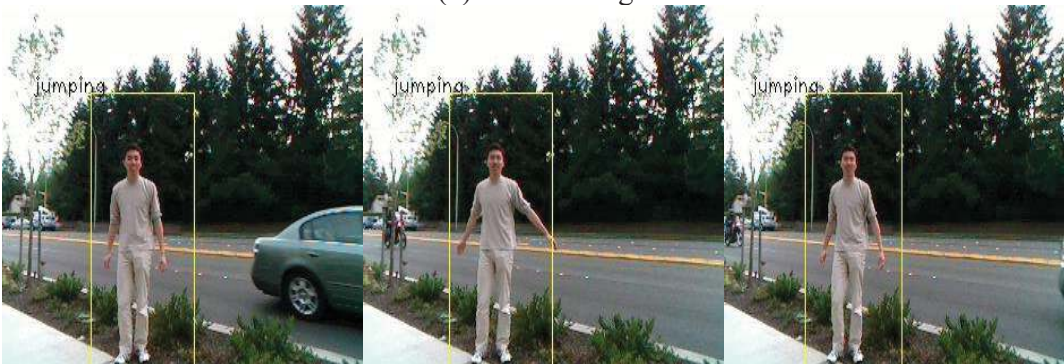
From Fig. 6.9, it can be observed that the person is performing “standing” activity at different viewing angles. From Fig.6.9, it is also clear that the proposed method is well capable of recognizing static and dynamic activities. Moreover, there is some little movement in each activity, i.e. pose of human object does not remain still for all the time.



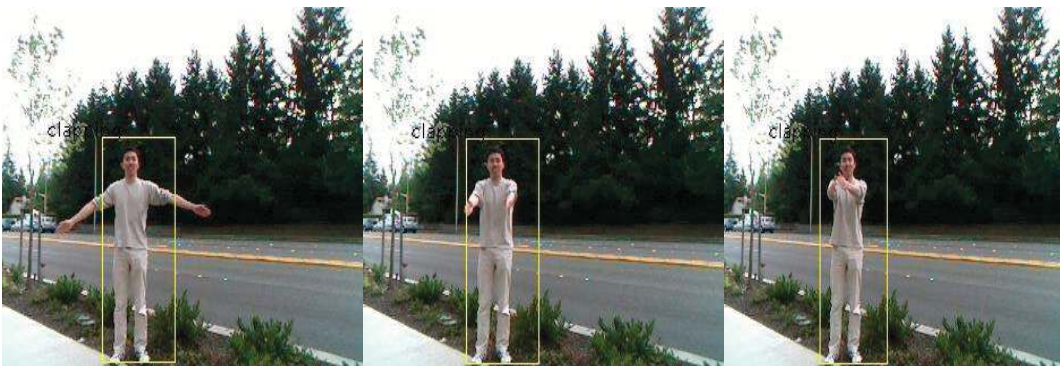
(a) Standing



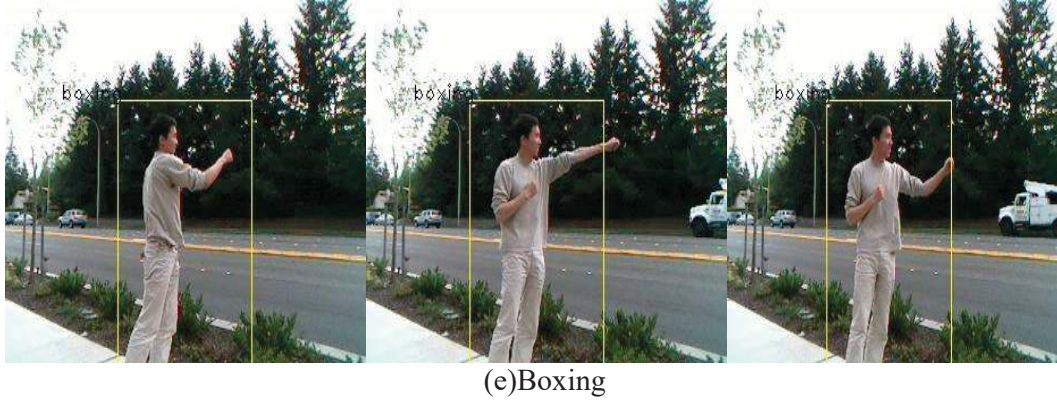
(b) Handwaving



(c) Jumping



(d) Hand-clapping



**Figure 6.9:** Recognition of Activities with MSR action recognition database [165] (a) Standing (b) Hand-waving (c) Jumping (d) Hand-clapping (e) Boxing

Direction of each human object also changes in different frames. Therefore, the proposed method is pose invariant and frontal view is not necessary for recognition of objects and suits for recognition of objects with frontal as well as side view. Hence, one can get correct visual results by using the proposed method. It is capable of recognizing the activity at these different viewing angles correctly and the proposed method is robust towards different rotations of the activity.

These confusion matrices and recognition results in Tables 6.9 & 6.10 indicate that the proposed method performs better than other methods.

**Table 6.9:** Confusion matrix for the proposed method and other methods

| Recognized Instances<br>→<br>Total Instances↓ | Standing | Handwaving | Jumping | Handclapping | Boxing |
|-----------------------------------------------|----------|------------|---------|--------------|--------|
| <b>Proposed method</b>                        |          |            |         |              |        |
| Standing                                      | 1        | 0          | 0       | 0            | 0      |
| Handwaving                                    | 0        | 1          | 0       | 0            | 0      |
| Jumping                                       | 0        | 0          | 1       | 0            | 0      |
| Handclapping                                  | 0        | 0          | 0       | 1            | 0      |
| Boxing                                        | 0        | 0          | 0       | 0            | 1      |
| <b>Chaaroui et al. [176]</b>                  |          |            |         |              |        |
| Standing                                      | 0.81     | 0.06       | 0.08    | 0.05         | 0      |
| Handwaving                                    | 0.10     | 0.84       | 0.06    | 0            | 0      |
| Jumping                                       | 0.14     | 0.06       | 0.78    | 0.01         | 0.01   |
| Handclapping                                  | 0.10     | 0.10       | 0.04    | 0.76         | 0      |
| Boxing                                        | 0        | 0.11       | 0       | 0.08         | 0.81   |
| <b>Bobick et al. [45]</b>                     |          |            |         |              |        |
| Standing                                      | 0.75     | 0.10       | 0.05    | 0.05         | 0.05   |
| Handwaving                                    | 0.10     | 0.71       | 0.10    | 0.05         | 0.04   |
| Jumping                                       | 0.14     | 0.11       | 0.73    | 0.02         | 0      |
| Handclapping                                  | 0.10     | 0.10       | 0.05    | 0.70         | 0.05   |
| Boxing                                        | 0.06     | 0.04       | 0.12    | 0            | 0.78   |
| <b>Ahmad et al [52]</b>                       |          |            |         |              |        |
| Standing                                      | 0.88     | 0.06       | 0       | 0.04         | 0.02   |
| Handwaving                                    | 0.03     | 0.95       | 0.02    | 0            | 0      |
| Jumping                                       | 0.05     | 0.03       | 0.90    | 0.02         | 0      |
| Handclapping                                  | 0.01     | 0.01       | 0.02    | 0.96         | 0      |
| Boxing                                        | 0.02     | 0.03       | 0       | 0.01         | 0.94   |
| <b>Iosifidis et al. [177]</b>                 |          |            |         |              |        |
| Standing                                      | 0.90     | 0.10       | 0       | 0            | 0      |
| Handwaving                                    | 0.07     | 0.83       | 0.10    | 0            | 0      |
| Jumping                                       | 0        | 0.14       | 0.86    | 0            | 0      |
| Handclapping                                  | 0        | 0.06       | 0.10    | 0.84         | 0      |
| Boxing                                        | 0.16     | 0          | 0       | 0            | 0.84   |
| <b>Weinland et al. [58]</b>                   |          |            |         |              |        |
| Standing                                      | 0.80     | 0.20       | 0       | 0            | 0      |
| Handwaving                                    | 0        | 0.82       | 0       | 0.18         | 0      |
| Jumping                                       | 0        | 0          | 0.81    | 0            | 0.19   |
| Handclapping                                  | 0        | 0          | 0       | 0.80         | 0.20   |
| Boxing                                        | 0        | 0          | 0.20    | 0            | 0.80   |
| <b>Iosifidis et al. [174]</b>                 |          |            |         |              |        |
| Standing                                      | 0.81     | 0          | 0.19    | 0            | 0      |
| Handwaving                                    | 0        | 0.85       | 0       | 0.15         | 0      |
| Jumping                                       | 0        | 0.10       | 0.90    | 0            | 0      |
| Handclapping                                  | 0        | 0.21       | 0       | 0.79         | 0      |
| Boxing                                        | 0.15     | 0          | 0       | 0            | 0.85   |

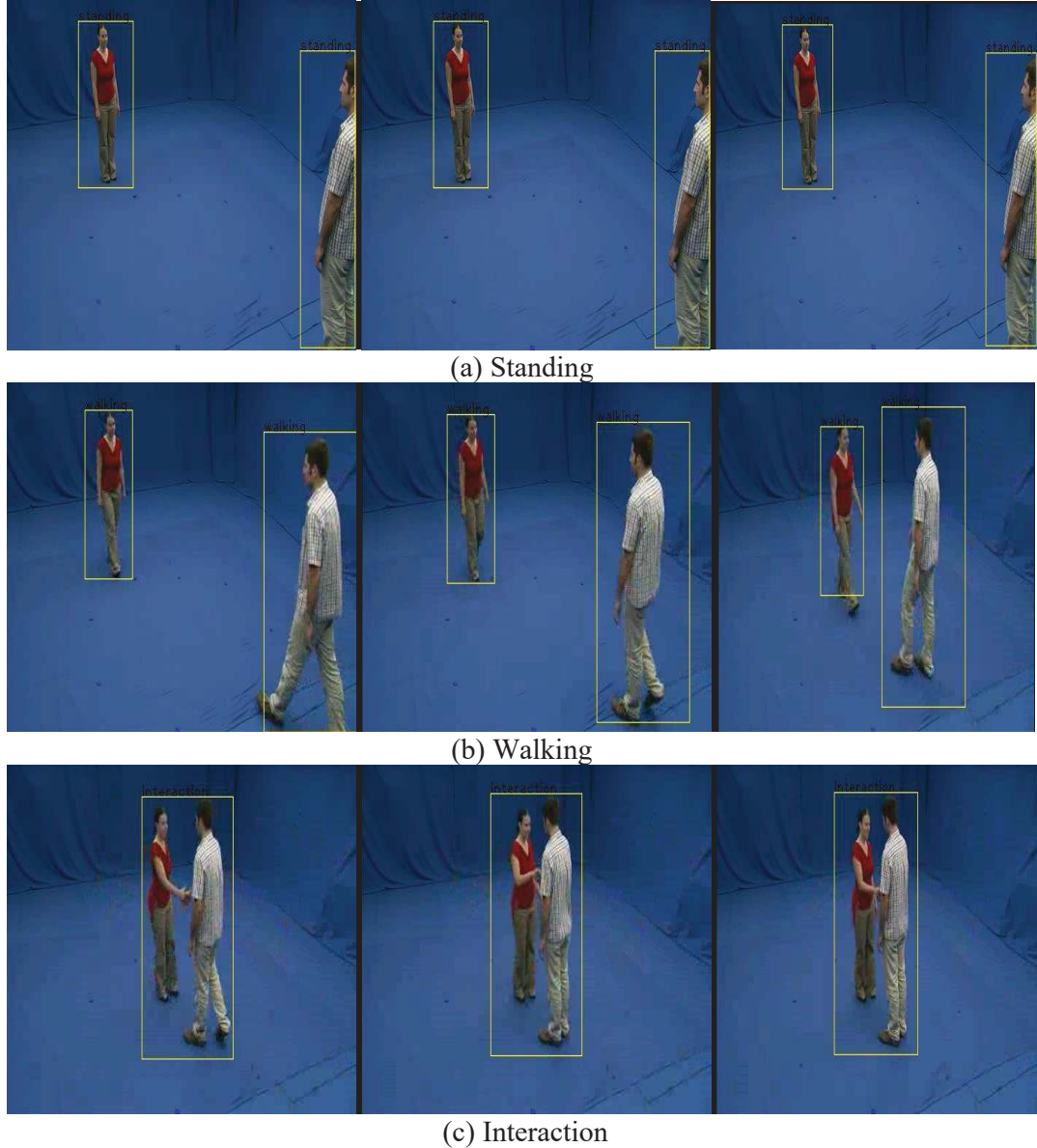
| <b>Weinland <i>et al.</i> [173]</b>  |      |      |      |      |      |
|--------------------------------------|------|------|------|------|------|
| <b>Standing</b>                      | 0.80 | 0.20 | 0    | 0    | 0    |
| <b>Handwaving</b>                    | 0    | 0.78 | 0    | 0.10 | 0.12 |
| <b>Jumping</b>                       | 0.17 | 0    | 0.83 | 0    | 0    |
| <b>Handclapping</b>                  | 0.08 | 0.10 | 0    | 0.82 | 0    |
| <b>Boxing</b>                        | 0    | 0    | 0.20 | 0    | 0.80 |
| <b>Iosifidis <i>et al.</i> [41]</b>  |      |      |      |      |      |
| <b>Standing</b>                      | 0.91 | 0    | 0.09 | 0    | 0    |
| <b>Handwaving</b>                    | 0    | 0.90 | 0    | 0.06 | 0.04 |
| <b>Jumping</b>                       | 0    | 0.05 | 0.91 | 0    | 0.04 |
| <b>Handclapping</b>                  | 0    | 0    | 0.12 | 0.88 | 0    |
| <b>Boxing</b>                        | 0.10 | 0    | 0    | 0    | 0.90 |
| <b>Iosifidis <i>et al.</i> [175]</b> |      |      |      |      |      |
| <b>Standing</b>                      | 0.95 | 0    | 0    | 0.05 | 0    |
| <b>Handwaving</b>                    | 0.08 | 0.92 | 0    | 0    | 0    |
| <b>Jumping</b>                       | 0    | 0    | 0.93 | 0    | 0.07 |
| <b>Handclapping</b>                  | 0    | 0    | 0.09 | 0.91 | 0    |
| <b>Boxing</b>                        | 0.10 | 0    | 0    | 0    | 0.90 |

**Table 6.10:** Recognition results over the MSR view-point action dataset [165]

| <b>Method</b>                 | <b>Accuracy (%)</b> |
|-------------------------------|---------------------|
| Proposed Method               | 100                 |
| Chaaroui <i>et al.</i> [176]  | 80                  |
| Bobick <i>et al.</i> [45]     | 73.4                |
| Ahmad <i>et al.</i> [52]      | 92.6                |
| Iosifidis <i>et al.</i> [177] | 85.4                |
| Weinland <i>et al.</i> [58]   | 80.6                |
| Iosifidis <i>et al.</i> [174] | 84                  |
| Weinland <i>et al.</i> [173]  | 80.6                |
| Iosifidis <i>et al.</i> [41]  | 90                  |
| Iosifidis <i>et al.</i> [175] | 92.2                |

## Experiment 6

In Fig.6.10, we have shown multiple human activity recognition using the i3DPost dataset which is a multi-view dataset [164]. This database includes different multiple human actions (such as standing, walking, interaction) taken with a multi-view camera set-up. For this database, the proposed method and the method by Iosifidis *et al.* [41] performs as per the requirements but other methods such as Chaaroui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], and Iosifidis *et al.* [175] do not recognize multiple human activities.



**Figure 6.10:** Recognition of Multiple Activities in i3DPost multi-view dataset [164] (a) Standing (b) Walking (c) Interaction

Now, quantitative results have been shown for i3DPost dataset [164] in Tables 6.11 and 6.12. From the confusion matrices and recognition results shown in Tables 6.11, one can find that the accuracy of the proposed method is better than Iosifidis *et al.* [41]. But other methods proposed by Chaaraoui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], and Iosifidis *et al.* [175] are not working for a multiple human activity dataset. Each confusion matrix shows the performance of a particular method for the i3DPost dataset. Diagonal

values in Table 6.11 indicate correct recognition rate for this purpose which are far better in case of the proposed method than the other methods. Comparison of recognition accuracy of Iosifidis *et al.* [41] with the proposed method has been shown in Table 6.12 (calculated by using Eq. 6.21). It shows that the performance of the proposed method is better than Iosifidis *et al.* [41].

**Table 6.11:** Confusion matrix for the proposed method and Iosifidis *et al.* [41]

| Recognized Instances<br>→<br>Total Instances ↓ | Standing | Walking | Interaction |
|------------------------------------------------|----------|---------|-------------|
| <b>Proposed method</b>                         |          |         |             |
| <b>Standing</b>                                | 0.98     | 0       | 0           |
| <b>Walking</b>                                 | 0        | 0.99    | 0           |
| <b>Interaction</b>                             | 0        | 0       | 0.98        |
| <b>Iosifidis <i>et al.</i> [41]</b>            |          |         |             |
| <b>Standing</b>                                | 0.92     | 0.04    | 0.04        |
| <b>Walking</b>                                 | 0.05     | 0.95    | 0           |
| <b>Interaction</b>                             | 0.02     | 0       | 0.98        |

**Table 6.12:** Recognition results over the i3DPost dataset [27]

| Method                       | Accuracy (%) |
|------------------------------|--------------|
| Proposed Method              | 98.33        |
| Iosifidis <i>et al.</i> [41] | 95           |

In Table 6.13, average computation time (second/frame) and memory consumption for different methods for a video of frame size 480 x 320 with 100 frames [165] are shown. From the Table 6.13, it can be observed that the proposed method is faster than Chaaraoui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], Iosifidis *et al.* [41], and Iosifidis *et al.* [175].

**Table 6.13:** Computational Time and Consumption Memory for MSR view-point action dataset [165]

| S.no. | Methods                       | Computational Time<br>(in frame/second) | Memory Consumption<br>(MB) |
|-------|-------------------------------|-----------------------------------------|----------------------------|
| 1     | Proposed Method               | <b>1.232</b>                            | <b>3.90</b>                |
| 2     | Chaaroui <i>et al.</i> [176]  | 1.722                                   | 22.92                      |
| 3     | Bobick <i>et al.</i> [45]     | 1.443                                   | 9.40                       |
| 4     | Ahmad <i>et al.</i> [52]      | 1.376                                   | 24.95                      |
| 5     | Iosifidis <i>et al.</i> [177] | 1.912                                   | 8.64                       |
| 6     | Weinland <i>et al.</i> [58]   | 1.753                                   | 7.08                       |
| 7     | Iosifidis <i>et al.</i> [174] | 1.325                                   | 11.37                      |
| 8     | Weinland <i>et al.</i> [173]  | 1.687                                   | 13.35                      |
| 9     | Iosifidis <i>et al.</i> [41]  | 1.912                                   | 30.92                      |
| 10    | Iosifidis <i>et al.</i> [175] | 1.412                                   | 17.62                      |

Also from the Table 6.13, the proposed method consumes only 3.90 megabytes of RAM which is the least in comparison with the other methods discussed [176, 45, 52, 177, 58, 174, 173, 41, 175]. Therefore, it can be concluded that the time required for the execution of the proposed method is faster than other methods and also consumes less amount of the system memory.

#### 6.4. Conclusions

In this chapter, we have proposed a multi-view human activity recognition system. This system is based on three consecutive modules. These modules are (i) background subtraction (ii) feature extraction and (iii) classification. The background subtraction was performed using change detection method. The contour-based pose features from silhouettes find the different key poses for human activities such as bending, standing, sitting etc. After that uniform rotation invariant LBP descriptor was computed. Its rotation invariant nature provides view invariant recognition for multi-view human activities. Multiclass SVM classifier has been applied for recognition of different activities. This approach was performed on six multi-view human activity video datasets: our own view point dataset, VideoWeb Multi-view dataset [166], i3DPost multi-view dataset [164],

WVU multi-view human action recognition dataset [167], MSR action recognition database [165] and i3DPost multi-view dataset [164] for multiple human. Qualitative and quantitative experimental results demonstrate the robustness of the proposed method against different viewpoints. The proposed method was compared with Chaaraoui *et al.* [176], Bobick *et al.* [45], Ahmad *et al.* [52], Iosifidis *et al.* [177], Weinland *et al.* [58], Iosifidis *et al.* [174], Weinland *et al.* [173], Iosifidis *et al.* [41], and Iosifidis *et al.* [175] and its performance found better than them.