



Technique for person re-identification using deep ranking

¹Vijayalaxmi Mudnoor, ² Sri. Naveen Kumar B,

¹PG student, Dept. of CSE, UBDTC, Davanagere

² Assistant professor Dept. of CSE, UBDTCE, Davanagere

Abstract— The system present a coherent, discriminative framework for simultaneously tracking multiple people and estimating their collective activities. Instead of treating the two problems separately, our model is grounded in the intuition that a strong correlation exists between a person's motion, their activity, and the motion and activities of other nearby people. we introduce a hierarchy of activity types that creates a natural progression that leads from a specific person's motion to the activity of the group as a whole.

Keywords- face recognition, context features, HMRF.

I. INTRODUCTION

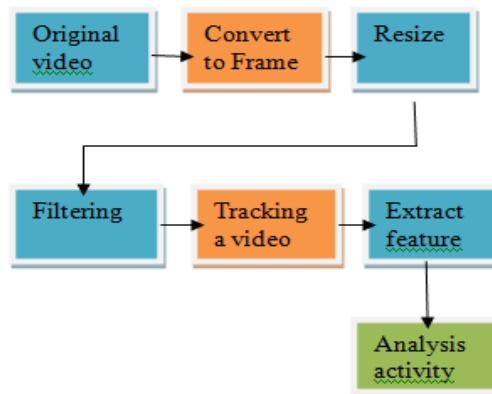
Surveillance videos in unconstrained environments typically consist of long duration sequences of activities which occur at different spatio-temporal locations and can involve multiple people acting simultaneously. Often, the activities have contextual relationships with one another. Although context has been studied in the past for the purpose of activity recognition, the use of context in recognition of activities in such challenging environments is relatively unexplored. In this paper, we propose a novel method for capturing the spatio-temporal context between activities in a Markov random field. Given a collection of videos and a set of weak classifiers for individual activities, the spatio-temporal relationships between activities are represented as probabilistic edge weights in the MRF. This model provides a generic representation for an activity sequence that can extend to any number of objects and interactions in a video.

II. PROPOSED SYSTEM

In proposed system the system propose a two-level hierarchical graphical model, which learns the relationship between tracks, and their corresponding activity segments, as well as the spatiotemporal relationships across activity segments. The HMRF is constructed on these track lets and activity segments. Using the obtained labels from recognition, the cost matrix is updated and the tracks are re-computed.

As a unifying framework to integrate the low- and high- level representations of human activity in video, we propose a hierarchical graphical model for recognizing human activities. In graphical models, the variables of interest are represented as the nodes and the relations between the variables are represented as links (or edges) that connect the corresponding nodes. Basically, graphical models (or graphs) are classified into two classes: directed graphs and undirected graphs. Directed graphs contain directed links that represent cause-effect relations between the nodes; a directed link denoted by an arrow originates from a cause variable and is directed toward an effect variable.

III. IMPLEMENTATION



- Frame conversion
- Filtering
- Track the video
- Feature extraction
- Analysis the activity

Frame conversion:

- Pre-processing consists of computing tracklets and computing low level features such as space-time interest points in the region around these tracklets.
- Tracking involves association of one or more tracklets to tracks.
- Activity localization can now be defined as a grouping of tracklets into activity segments and recognition can be defined as the task of labelling these activity segments.

Filtering:

- The system assumes that we have with us a set of tracklets, which are short-term fragments of tracks with low probability of error.
- Tracklets have to be joined to form long-term tracks. In a multi-person scene, this involves tracklet association. Here, we use a basic particle filter for computing tracklets as mentioned.
- For the test video, it is assumed that each tracklet belongs to a single activity.

Track the video:

- To begin with, we generate a set of match hypotheses for tracklet association and a likely set of tracks.
- An observation potential is computed for each tracklet using the features computed at the tracklet.
- Tracklets are grouped into activity segments using a standard baseline classifier such as multiclass SVM or motion segmentation.

Feature extraction:

- The node features and edge features for the potential functions are computed from the training data.
- There are two tasks to be performed on the graph - choosing an appropriate structure and learning the parameters of the graph. Both these steps can be performed simultaneously by posing the parameter learning as an L1-regularized optimization.

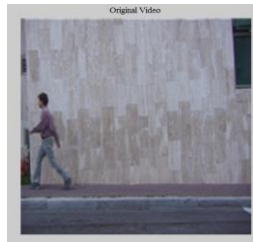
Analysis the activity:

- To evaluate the accuracy of activity recognition, if there is more than a 40% overlap in the spatiotemporal region of a detected activity as compared to the ground truth and the labelling corresponds to the ground truth labelling, the recognition is assumed to be correct.
- It can be used in conjunction with many other types of learning algorithms to improve their performance.

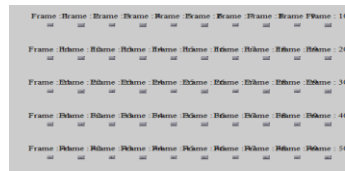
- The output of the other learning algorithms ('weak learners') is combined into a weighted sum that represents the final output of the boosted classifier.

IV. RESULTS

1.ORIGINAL VIDEO



2 VIDEOS CONVERT TO FRAMES



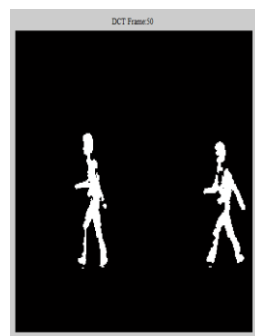
3.RESIZED IMAGE



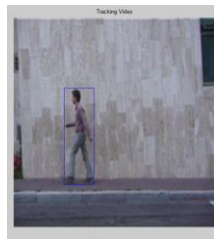
4.BINARY IMAGES



5.DCT FRAMES



6. TRACKING VIDEOS



7. OBJECT IN FRAMES



V. CONCLUSION

Spatio temporal contextual relationships between activities and the influence of tracks on them has been modelled using the graph. The activity labels obtained in the bottom-up processing are in turn used to correct the errors in tracking in a top-down approach. In the interaction hierarchy, two-person interaction is defined in terms of the combination of two single-person actions. We have developed a dynamic Bayesian network to combine the poses into gestures, and a rule-based decision tree to classify human interactions occurring between two persons. The performance of the system depends on several factors. The high-level processing relies on the robustness of the low level processing.

REFERENCES

- [1] M. R. Amer and S. Todorovic, "Sum-product networks for modeling activities with stochastic structure," in Proc. IEEE Conf. Compute. Vis. Pattern Recognit., Jun. 2012, pp. 1314–1321.
- [2] W. Brendel and S. Todorovic, "Learning spatiotemporal graphs of human activities," in Proc. IEEE Int. Conf. Compute. Vis., Nov. 2011, pp. 778–785.
- [3] V. Chandrasekaran, N. Srebro, and P. Harsha, "Complexity of inference in graphical models," in Proc. 24th Annu. Conf. Uncertainty Artif. Intell., 2008, pp. 70–78.
- [4] C.-Y. Chen and K. Grauman, "Efficient activity detection with max-sub graph search," in Proc. IEEE Conf. Compute. Vis. Pattern Recognit., Jun. 2012, pp. 1274–1281.