

Histograms of Oriented Gradients for Live Capturing and Tracking of Humans

Kiran Patil¹, Dr A J Patil²

¹PG Scholar, Dept of Electronics and Telecommunication, Gulabrao Deokar, COE, Jalgaon, India

²Principal, Gulabrao Deokar, COE, Jalgaon, India

ABSTRACT

The histogram of oriented gradients descriptor is one of the best and most popular descriptors used for pedestrian detection. Unfortunately this technique suffers from one big problem: speed. Like most sliding window algorithms it is very slow, making it unsuitable for any real time application. The HOG detector is a sliding window algorithm. This means that for any given image a window is moved across at all locations and scales and a descriptor is computed. For that window a pre trained classifier is used to assign a matching score to the descriptor. The classifier used is a linear SVM classifier and the descriptor is based on histograms of gradient orientations. Gradient orientations and magnitude are obtained for each pixel from the pre-processed image. If a color image is used the gradient with the maximum magnitude (and its corresponding orientation) is chosen. This is done by convoluting the image with the 1-D centred kernel $[-1 \ 0 \ 1]$ by rows and columns. Several other techniques have been tried (including 3x3 Sobel mask or 2x2 diagonal masks) but the simple 1-D centred kernel gave the best performance (for pedestrian detection). The first is the idea of using appearance information, whose importance in object detection has been widely demonstrated, and specially the idea of combining cues with the HOG descriptor, e.g., co-occurrence HOG color HOG etc. It has been largely demonstrated by the proposal of different descriptors that appearance is an important cue for object detection. The second source is the increasing trend of using segmentation for both detection and segmentation, which in our case has the potential of highlighting the shape of the human.

Keyword: HOG, SVM, image segmentation, video segmentation.

1. INTRODUCTION

This project is implemented for the domain specific human detection and the image retrieval in video by exploring the techniques for feature extraction, semantic annotation and query processing. Feature extraction and semantic annotation are the main requirements to show that the indexing method is feasible; that is the information can be extracted without too much manual intervention. Query processing that includes developing a data model, choosing a suitable query language and constructing a compact representation, is the most important aspect to show the success of the indexing method in terms of satisfying users/applications' requirements.

Detecting humans in images is a challenging task owing to their variable appearance and the wide range of poses that they can adopt. The first need is a robust feature set that allows the human form to be discriminated cleanly, even in cluttered backgrounds under difficult illumination. We study the issue of feature sets for human detection, showing that locally normalized Histogram of Oriented Gradient (HOG) descriptors provide excellent performance relative to other existing feature sets including wavelets. The proposed Descriptors are reminiscent of edge orientation histograms, SIFT descriptors and shape contexts, but they are computed on a dense grid of uniformly spaced cells and they use overlapping local contrast normalizations for improved performance. We make a detailed study of the effects of various implementation choices on detector performance, taking pedestrian detection" (the detection of mostly visible people in more or less upright poses) as a test case. For simplicity and speed, we use linear SVM as a baseline classifier throughout the study. The new detectors give essentially perfect results on the MIT pedestrian test set, so we have created a more

challenging set containing over 1800 pedestrian images with a large range of poses and backgrounds. On-going work suggests that our feature set performs equally well for other shape-based object classes

2. HOG DESCRIPTOR AND USAGE FOR OBJECT DETECTION

The histogram of oriented algorithm for object detection was introduced in [1]. An example of detection results for pedestrians and heads is visible in Figure 1.



Figure 1: Example results after applying the HOG detector for pedestrians (left) and heads (right)

The HOG detector is a sliding window algorithm. This means that for any given image a window is moved across at all locations and scales and a descriptor is computed. For that window a pre trained classifier is used to assign a matching score to the descriptor. The classifier used is a linear SVM classifier and the descriptor is based on histograms of gradient orientations. Firstly the image is padded and gamma normalization is applied. Padding means that extra rows and columns of pixels are added to the image. Each new pixel gets the color of its closest pixel from the original image. This helps the algorithm deal with the case when a person is not fully contained inside the image. The gamma normalization has been proven to improve performance for pedestrian detection but it may decrease performance for other object classes. To compute the gamma correction the color for each channel is replaced by its square root.

Gradient orientations and magnitude are obtained for each pixel from the pre-processed image. If a color image is used the gradient with the maximum magnitude (and its corresponding orientation) is chosen. This is done by convoluting the image with the 1-D centered kernel $[-1 \ 0 \ 1]$ by rows and columns. Several other techniques have been tried (including 3x3 Sobel mask or 2x2 diagonal masks) but the simple 1-D centered kernel gave the best performance (for pedestrian detection).

Each detection window is divided into several groups of pixels, called cells. For each cell a histogram of gradient orientations is computed. For pedestrian detection the cells are rectangular and the histogram bins are evenly spaced between 0 and 180 degrees. The contribution of each pixel to the cell histogram is weighted by a Gaussian centered in the middle of the block and then interpolated tri linearly in orientation and position. This has the effect of reducing aliasing.

3. HOG RE-WEIGHTING USING GLOBAL IMAGE SEGMENTATION

The proposed approach consists in re-weighting the HOG descriptor for each one of the cells while it is computed. Basically, HOG consists in an intelligent grouping of gradient information (cells and blocks), as well as well-engineered histograms of gradient orientations (weighting by gradient magnitude, bin interpolation, histogram normalization and outliers clipping are the major steps 1). The window of interest is covered by overlapping blocks; therefore, the HOG descriptor of the whole window usually ends up being thousand-dimensional. A linear SVM is used for learning the human classifier works in such high dimensional space. Accordingly, the obtained classifier is just a weighted summation running on such number of dimensions. When

computing HOG, the gradient orientation θ_P at a given pixel P is weighted by the corresponding magnitude μ_P , i.e., μ_P is accumulated in the histogram bin corresponding to θ_P . Notice that the gradient at P , only encodes local differences in intensity or color, i.e., differences between adjacent pixels. In this paper, we want to incorporate differences based on a wider spatial support into the process in order to assess if they allow obtaining a human detector with higher performance. In particular, we want to weight μ_P by a given ω_P coming up from an image segmentation process, i.e., the vote of θ_P in the histogram will be the re-weighted magnitude λ_P instead of μ_P . Fig. 2 illustrates the idea.

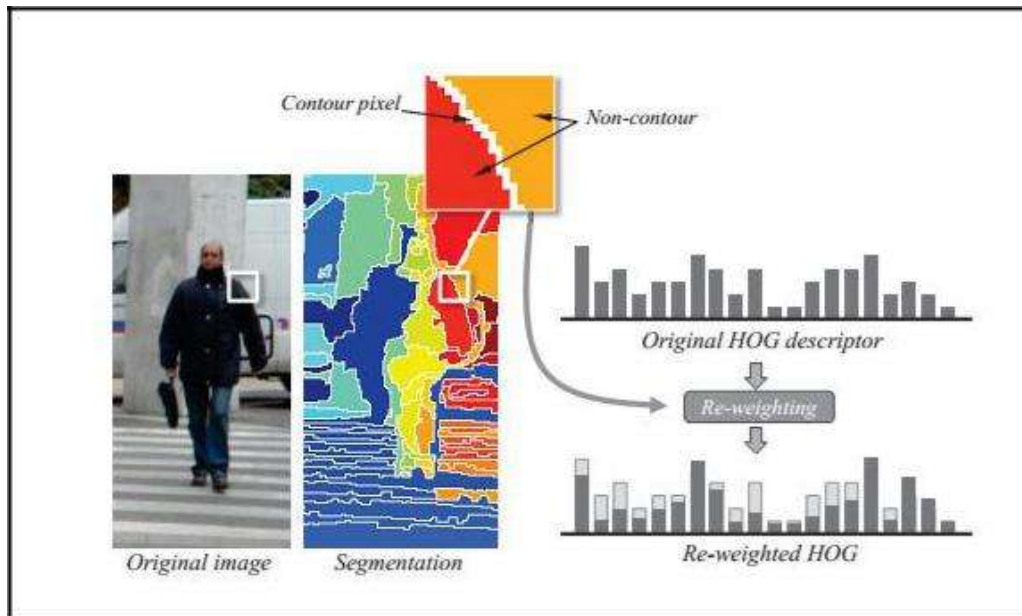


Fig. 2: Re-weighting of the HOG descriptor according to the image segmentation cues.

Here, in order to compute the weighted HOG, we apply the selected image segmentation algorithm (i.e., S) to each level of the pyramid. Thus, we obtain a sort of multi-scale segmentation of the original image (Fig. 3).

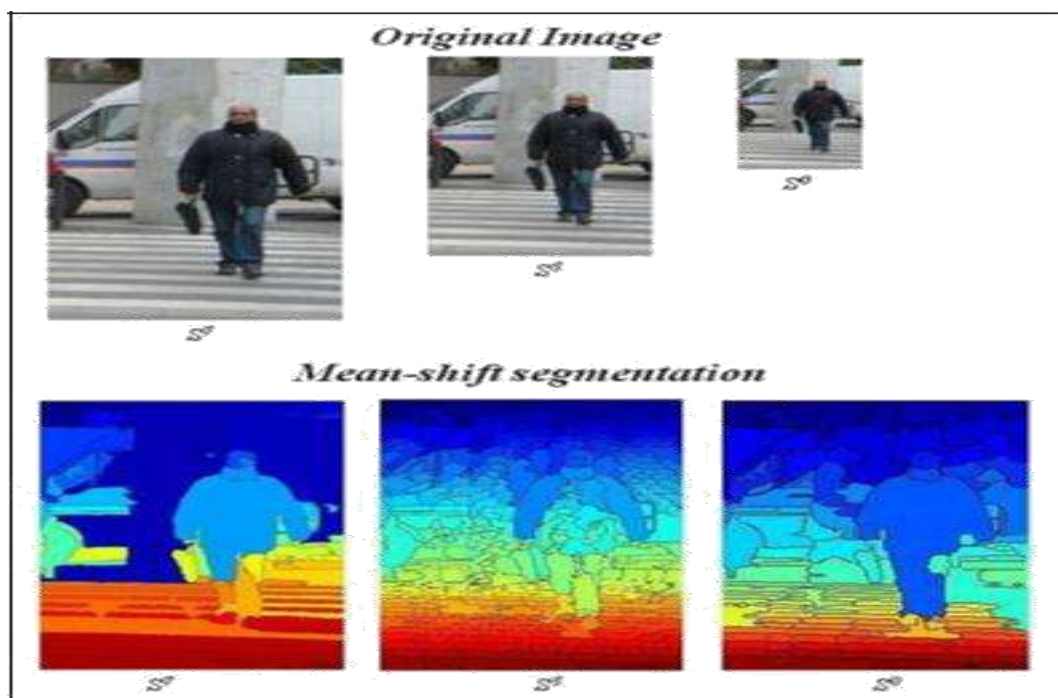


Fig. 3: Pyramid-segmentation. The scale of the slice in the pyramid affects the segmentation, similarly as it affects the HOG descriptor.

4. PROPOSED WORK

The proposed work is divided into two parts

1. The first part detects and counts the total number of visitors entering in specified domain.
2. The second part of the proposed work extracts an image from an video and retrieves more videos related to the image using FP growth and SVM.

This section gives an overview of our feature extraction chain, which is summarized in Fig. 1. Implementation details are postponed until x6. The method is based on evaluating well-normalized local histograms of image gradient orientations in a dense grid. Similar features have seen increasing use over the past decade. The basic idea is that local object appearance and shape can often be characterized rather well by the distribution of local intensity gradients or edge directions, even without precise knowledge of the corresponding gradient or edge positions. In practice this is implemented by dividing the image window into small spatial regions ("cells"), for each cell accumulating a local 1-D histogram of gradient directions or edge orientations over the pixels of the cell. The combined histogram entries form the representation. For better in variance to illumination, shadowing, etc., it is also useful to contrast-normalize the local responses before using them. This can be done by accumulating a measure of local histogram energy" over somewhat larger spatial regions ("blocks") and using the results to normalize all of the cells in the block. We will refer to the normalized descriptor blocks as Histogram of Oriented Gradient (HOG) descriptors. Tiling the detection window with a dense (in fact, overlapping) grid of HOG descriptors and using the combined feature vector in a conventional SVM based window classifier gives our human detection chain in Fig 4.

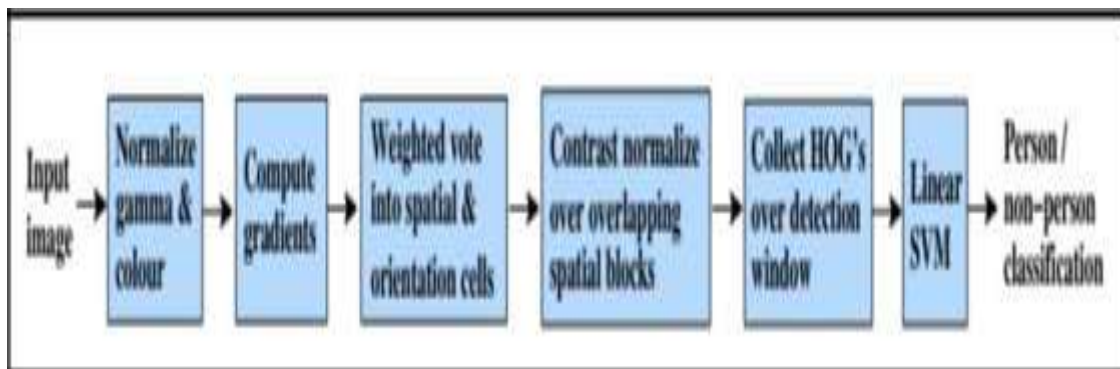


Fig 4. Human detection Chain

The use of orientation histograms has many precursors, but it only reached maturity when combined with local spatial histogramming and normalization in Lowe's Scale Invariant Feature Transformation (SIFT) approach to wide baseline image matching, in which it provides the underlying image patch descriptor for matching scale in variant key-points. SIFT-style approaches perform remarkably well in this application. The Shape Context work studied alternative cell and block shapes, albeit initially using only edge pixel counts without the orientation histogramming that makes the representation so effective.

The success of these sparse feature based representations has somewhat overshadowed the power and simplicity of HOG's as dense image descriptors. We hope that our study will help to rectify this. In particular, our informal experiments suggest that even the best current key point based approaches are likely to have false positive rates at least 1–2 orders of magnitude higher than our dense grid approach for human detection, mainly because none of the key point detectors that we are aware of detect human body structures reliably.

The HOG/SIFT representation has several advantages. It captures edge or gradient structure that is very characteristic of local shape and it does so in a local representation with an easily controllable degree of invariance to local geometric and photometric transformations: translations or rotations make little difference if they are much smaller than the local spatial or orientation bin size. For human detection, rather coarse spatial sampling, fine orientation sampling and strong local photometric normalization turns out to be the best strategy, presumably because it permits limbs and body segments to change appearance and move from side to side quite a lot provided that they maintain a roughly upright orientation.

CLUSTERING is a natural solution to abbreviate and organize the content of a video. A preview of the video content can simply be generated by showing a subset of clusters or the representative frames of each

cluster. Similarly, retrieval can be performed in an efficient way since similar shots are indexed under the same cluster. Regardless of these advantages, general solutions for clustering video data are a hard problem. For certain applications, motion features dominate the clustering results; for others, visual cues such as color and texture are more important. Moreover, for certain types of applications, decoupling of camera and object motions ought to be done prior to clustering.

Video retrieval techniques, to date, are mostly extended directly or indirectly from image retrieval techniques. Examples include first selecting key frames from shots and then extracting image features such as color and texture features from those key frames for indexing and retrieval. The success from such extension, however, is doubtful since the spatial-temporal relationship among video frames is not fully exploited. Due to this consideration, recently more works have been dedicated to address more specifically, to exploit and utilize the motion information along the temporal dimension for retrieval.

5. INSERTED CAPTION DETECTION

During or after a key event in most sport videos, a text with surrounding box is usually inserted to draw the attention of users to some content-sensitive information. For example, after a goal is scored in a soccer game, some texts are usually displayed to inform viewers about the current score-line. To detect text in video, Wernicke and Lienhart used the gradient of color image to calculate the complex-values from edge orientation image. It is defined to map all edge orientation between 0 and 90 degrees, thereby distinguishing horizontal, diagonal and vertical lines. Similarly, Mita and Hori localized character regions by extracting strong still edges and pixels with a stable intensity for two seconds.

Strong edges are detected by Sobel filtering and verified to be standstill by comparing four consecutive gradient images. Thus, current text detection techniques detect the edges of the rectangle box formed by text regions in color video-frames and then check if the edge will stay for more than 2 seconds.

Based on these concepts, the essence of our text display detection method is that sport videos use only horizontal text in 99% of the cases, as can be seen in Figure 5.5. If we can detect strong horizontal line in a frame, we can locate the starting point of a text region. The main steps involved in detecting text display detection can be seen in Figure 5.7 and are explained below.

First, video track is segmented into one minute clip. Each frame for every one second gap is pre-processed to optimize performance by converting the color scheme to gray scale and the size is reduced to a smaller pre-set value. Second, Sobel Filter is performed to calculate the edge (gradient) of the current frame and Hough Transform is used on the gradient image to detect line spaces (R) between 0 - 1800. Threshold1 is applied on the R values to detect the potential candidates of strong lines which are usually formed by the box surrounding text displays. After these lines are detected, the system calculates the r (rho) and t (theta) values from the peak coordinates. It should be noted that r value indicates the location of the line in terms of the number of pixels from the center and t indicates the angle of the line. In order to verify that the detected lines are the candidates of text display, the lines are filtered out using two criteria:

- 1) The absolute value of r is less than n % of the maximum y -axis, and
- 2) The corresponding t is equal to 900 (horizontal). This n -value is regarded as threshold2.

The first criterion is important to ensure that the location of the lines is within the usual location of text display. The second criterion is to ensure that the line is horizontal because some strong horizontal lines can be detected from other area beside the text display, such as the boundary between field and crowd. Finally, for each of the lines detected the system check that their location (i.e. the r values) is consistent for at least m seconds. This m -value is regarded as threshold3. The purpose of this check is to ensure that the location of the lines is consistent with the next frames as text displays always appear for at least 2 seconds to give viewers enough time to read. In fact, when a text display size is large and contains lots of information, it will be displayed even longer to give viewers enough time to read. Figure 2.16 illustrates how Sobel-Hough transform can be used to detect text.

Since there are not so many texts during a sport match (otherwise they can be distractive), it takes longer to check text display in all video frames. Instead, specific domain knowledge should be used to predict text appearances in sport videos which are quite typical from one match to another. Fig 5 shows an example of predicting text occurrences based on specific domain knowledge for soccer video. However, during the experiment to detect generic events, the system still check all frames since text display can detect some events which sometimes cannot be detected by whistle and excitement sounds. For example, whistle does not always exist during start of a match, but texts are usually displayed at the beginning of a match to show essential information like team formation.

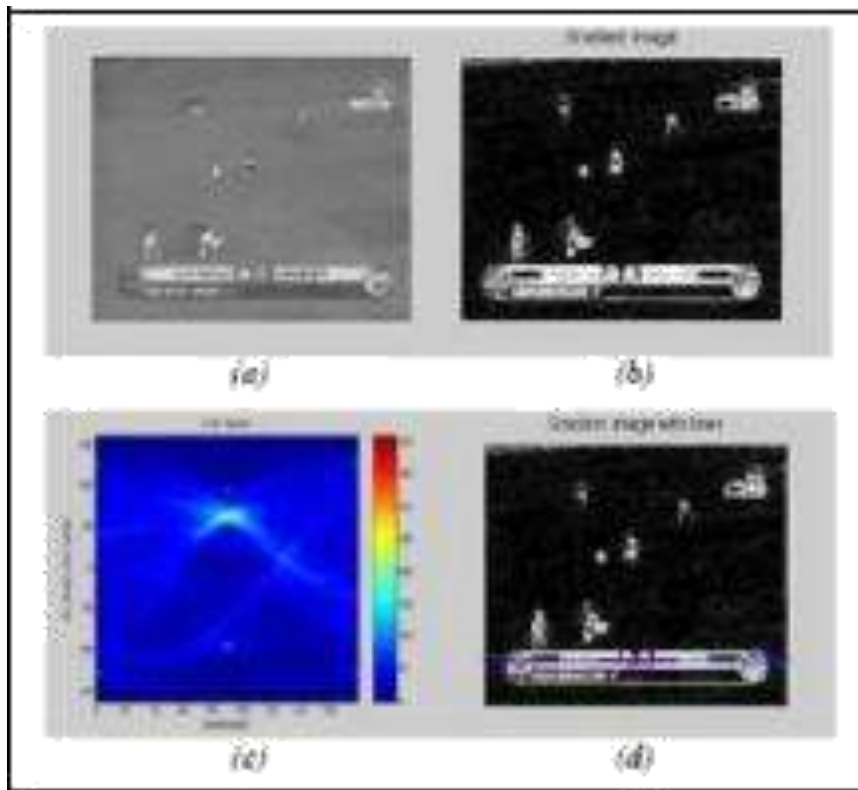


Fig 5. How Sobel/Hough Transform Detect Text: (a) Gray scaled& resized frame, (b) Gradient image reveals the lines (c) Peaks in Hough transform, (d) Horizontal lines detected

6. CONCLUSION

In this work I have to investigate the possibility of improving HOG descriptors in the context of human detection. In particular by weighting HOG with information coming from image segmentation. The detected images are all stored for mining to count the number of people in a specific domain. The specific events detection can be improved and extended by deriving some additional statistical features such as the location of slow motion replay of the last close-up frame, and including more mid-level features which can improve the highlights detection.

The work also retrieves the videos from net on the retrieved image on the video as video tracking.

7. References

1. Arbel' aez, P., Maire, M., Fowlkes, C., Malik, J.: Contour detection and hierarchical image segmentation. *IEEE Trans. on Pattern Analysis and Machine Intelligence*99(1), 898–916(2010).
2. Boix, X., Gonfaus, J., van de Weijer, J., Bagdanov, A., Serrat, J., Gonz` alez, J.: Harmony potentials for joint classification and segmentation. In: *IEEE Conf. on Computer Vision and Pattern Recognition*. San Francisco, CA, USA (2010).
3. Carreira, J., Sminchisescu, C.: Constrained parametric min-cuts for automatic object segmentation (2010).
4. Comaniciu, D., Meer, P.: Mean shift: a robust approach toward feature space analysis. *IEEE Trans. on Pattern Analysis and Machine Intelligence*24(5), 603–619 (2002).

5. Dalal, N.: Finding people in images and videos. PhD Thesis, Institut National Polytechnique de Grenoble / INRIA Rhône-Alpes (2006).
6. Dalal, N., Triggs, B.: Histograms of oriented gradients for human detection. In: IEEE Conf. on Computer Vision and Pattern Recognition. San Diego, CA, USA (2005).
7. F. Bergeaud and S. Mallat. Matching pursuit of images. In Image Processing, 1995. Proceedings., International Conference on, volume 1, pages 53–56. IEEE, 1995.
8. M. Everingham, L. Van Gool, C. K. I. Williams, J. Winn, and A. Zisserman. The pascal visual object classes (voc) challenge. International Journal of Computer Vision, 88(2):303–338, June 2010.
9. Z. Rahman, D. Jobson, and G. Woodell. Multi-scale retinex for color image enhancement. In Image Processing, 1996. Proceedings., International Conference on, volume 3, pages 1003–1006. IEEE, 1996.