

Aggregate Channel Features for Multi-view Face Detection

Bin Yang Junjie Yan Zhen Lei Stan Z. Li*

Center for Biometrics and Security Research & National Laboratory of Pattern Recognition
Institute of Automation, Chinese Academy of Sciences, China

yb.derek@gmail.com {jjyan, zlei, szli}@nlpr.ia.ac.cn

Abstract

Face detection has drawn much attention in recent decades since the seminal work by Viola and Jones. While many subsequences have improved the work with more powerful learning algorithms, the feature representation used for face detection still can't meet the demand for effectively and efficiently handling faces with large appearance variance in the wild. To solve this bottleneck, we borrow the concept of channel features to the face detection domain, which extends the image channel to diverse types like gradient magnitude and oriented gradient histograms and therefore encodes rich information in a simple form. We adopt a novel variant called aggregate channel features, make a full exploration of feature design, and discover a multi-scale version of features with better performance. To deal with poses of faces in the wild, we propose a multi-view detection approach featuring score re-ranking and detection adjustment. Following the learning pipelines in Viola-Jones framework, the multi-view face detector using aggregate channel features shows competitive performance against state-of-the-art algorithms on AFW and FDDB test-sets, while runs at 42 FPS on VGA images.

1. Introduction

Human face detection have long been one of the most fundamental problems in computer vision and human-computer interaction. In the past decade, the most influential work should be the face detection framework proposed by Viola and Jones [22]. The Viola-Jones (abbreviated as VJ below) framework uses rectangular Haar-like features and learns the hypothesis using Adaboost algorithm. Combined with the attentional cascade structure, the VJ detector achieved real-time face detection at that time. Despite the great success of the VJ detector, the performance is still far from satisfactory due to the large appearance variance of faces in unconstrained settings.

*Corresponding author.

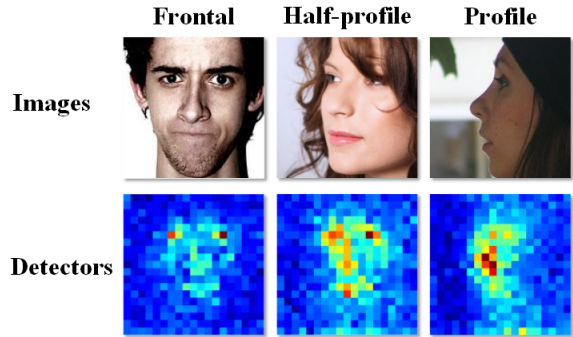


Figure 1. An intuitive visualization of our multi-view face detector using aggregate channel features. The area with warmer color indicates more attention paid to by the detector.

To handle faces in the wild, many subsequences of VJ framework merged. These methods mainly get the performance gains in two aspects, more complicated features [17, 19, 26] and (or) more powerful learning algorithms [14, 1, 25]. As the combination of boosting and cascade has been proven to be quite effective in face detection, the bottleneck lies in the feature representation since complicated features adopted in the above literatures bring about limited performance gains at the cost of large computation cost.

Lately in another domain of pedestrian detection, a family of channel features has achieved record performances [6, 5]. Channel features compute registered maps of the original images like gradients and histograms of oriented gradients and then extract features on these extended channels. The classifier learning process follows the VJ framework pipeline. In this paper, we adopt a variant of channel features called aggregate channel features [5], which are extracted directly as pixel values on subsampled channels. Channel extension offers rich representation capacity, while simple feature form guarantees fast computation. With these two superiorities, the aggregate channel features break through the bottleneck in VJ framework and have the potential to make great advance in face detection.

As we mainly concentrate our efforts to the feature representation rather than learning algorithms in this paper, we not only just adopt the aggregate channel features in face detection, but also try to explore the full potential of this novel representation. To do so, we make a deep and all-round investigation into the specific feature parameters concerning channel types, feature pool size, subsampling method, feature scale and so on, which gives insights into the feature design and hopefully provides helpful guidelines for practitioners. Through the deep exploration, we find that: 1) multi-scaling the feature representation further enriches the representation capacity since original aggregate channel features have uniform feature scale; 2) different combinations of channel types impact the performance greatly, while for face detection the color channel in LUV space, plus gradient magnitude channel and gradient histograms channels in RGB space show best result; 3) multi-view detection is proven to be a good match with aggregate channel features as the representation naturally encodes the facial structure (Figure 1).

Although multi-view detection could effectively deal with diverse poses, additional issues come up as how to merge detections output by separately trained subview detectors, and how to deal with the offsets of location and scale between output detections and ground-truth. We solve these problems by carefully designed post-processing including score re-ranking, detection merging and bounding box adjustment.

The detailed experimental exploration of aggregate channel features, along with our improvements on multi-view detection, leads to large performance gain in face detection in the wild. On two challenging face databases, AFW and FDDB, the proposed multi-view face detector shows competitive performance against state-of-the-art detectors in both detection accuracy and speed.

The remaining parts of this paper are organized as follows. Section 2 revisits related work in face detection. Section 3 describes how we build the face detector using aggregate channel features. Section 4 addresses problems concerning multi-view face detection. Experimental results on AFW and FDDB are shown in section 5 and we conclude the paper in section 6.

2. Related work

Face detection has drawn much attention since the early time of computer vision. Although many solutions had been put forward, it was not until Viola and Jones [22] proposed their milestone work that face detection saw surprising progress in the past decades. The VJ face detector features in three aspects: fast feature computation via integral image representation, classifier learning using Adaboost, and the attentional cascade structure. One main drawback of the VJ framework is that the features have limited repre-

sentation capacity, while the feature pool size is quite large to compensate for that. Typically, in a 24×24 detection window, the number of Haar-like features is 160,000 [22]. To address the problem, efforts are made in two directions. Some focus on more complicated features like HoG [26], SURF [13]. Some aim to speed up the feature selection in a heuristic way [18, 2]. However, the problem hasn't been solved perfectly. In this paper, we mainly focus on the feature representation part and make a deep exploration into it, which is complementary to existing work on the learning algorithm and classifier structure in the VJ framework.

Recently channel features have been proposed and shown record performance in pedestrian detection [6, 5]. Due to the channel extension to diverse types like gradients and local histograms, the features show richer representation capacity for classification. However, the features are extracted as rectangular sums at various locations and scales which we believe leads to a redundant feature pool. During preparation of this paper, Mathias *et al.* [16] independently discover the effectiveness of integral channel features in face detection domain. In this paper, we adopt a novel variant of channel features called aggregate channel features, which extract features directly as pixel values in extended channels without computing rectangular sums at various locations and scales. The feature has powerful representation capacity and the feature pool size is only several thousands. Through careful design in section 3 and implementation of multi-view detection in section 4, the aggregate channel features based detector achieves state-of-the-art performance on challenging databases.

3. Proposed face detector

In this section, we make a full exploration of the aggregate channel features in the context of face detection. We first give a brief introduction of the feature itself, including its computation, properties and advantages over traditional Haar-like features used in VJ framework. Then the detailed experimental investigation is described in two parts, feature design and training design. Before that, some guidelines concerning how we conduct the investigation are demonstrated. Each design part is divided into several separate experiments ended with a summary explaining the specific parameters used in our proposed face detector. Note that each experiment focuses on only one parameter and the others remain constant. Through the well-designed experiments, the proposed face detector based on aggregate channel features is built step by step. Issues concerning the implementation of multi-view face detection which further improves the performance are discussed in the next section.

3.1. Feature description

Channel extension: The basic structure of the aggregate channel features is channel. The application of channel

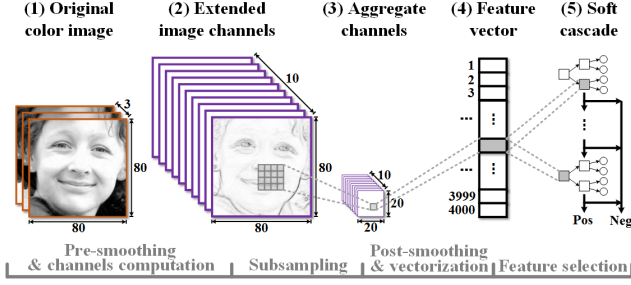


Figure 2. Work-flow of proposed face detector.

has a long history since digital images were invented. The most common type of channel should be the color channels of the image, with Gray-scale and RGB being typical ones. Besides color channels, many different channel types have been invented to encode different types of information for more difficult problems. Generally, channels can be defined as a registered map of the original image, whose pixels are computed from corresponding patches of original pixels [6]. Different channels can be computed with linear or non-linear transformation of the original image. To allow for sliding window detection, the transformations are constrained to be translationally invariant.

Feature computation: Based on the definition of channels, the computation of aggregate channel features is quite simple. As shown in Figure 2, given a color image, all defined channels are computed and subsampled by a pre-set factor. The aggregate pixels in all subsampled channels are then vectorized into a pixel look-up table. Note that an optional smoothing procedure can be done on each channel with a binomial filter both before computation and after subsampling.

Classifier learning: The learning process is quite simple. Two changes are made compared with VJ framework. First is that weak classifier is changed from decision stump to depth-2 decision tree. The more complex weak classifier shows stronger ability in seeking the discriminant intra and inter channel correlations for classification [15]. Second difference is that soft-cascade [1] structure is used. Unlike the attentional cascade structure in VJ framework which has several cascade stages, a single-stage classifier is trained on the whole training data and a threshold is then set after each weak classifier picked by Adaboost. These two changes lead to more efficient training and detection.

Overall superiority: Compared with traditional Haar-like features used in VJ framework, aggregate channel features have the following differences and advantages: 1) The image channels are extended to more types in order to encode diverse information like color, gradients, local histograms and so on, therefore possess richer representation capacity. 2) Features are extracted directly as pixel values on downsampled channels rather than computing rectangular

sums with various locations and scales using integral images, leading to a faster feature computation and smaller feature pool size for boosting learning. With the help of cascade structure, detection speed is accelerated more. 3) Due to its structure consistence with the overall image, when coupled with boosting method, the boosted classifier naturally encodes structured pattern information from large training data (see Figure 1 for an illustration), which gives more accurate localization of faces in the image.

3.2. Investigation guidelines

All investigations are trained on the AFLW face database¹ [10] and tested on the Annotated Faces in the Wild (AFW) testset². To make it clear, there are in total 36,112 positive samples and 108,336 negative samples selected from AFLW which are kept constant in all investigations. Testset contains 205 natural images with faces that vary a lot in pose, appearance and illumination.

To alleviate the ground-truth offset caused by different annotation styles (Figure 4) in training and testing set and make the evaluation more comparable, a lower Jaccard index³ with threshold 0.3 is adopted in comparative evaluation. Practically the lower threshold won't cause errors being mistakenly corrected. Note that in final evaluation of the proposed face detector (section 5), the AFW testset, together with another face benchmark Fddb database, are used as testbed and the evaluation metric follows the database protocol.

3.3. Feature design

To fully exploit the power of aggregate channel features in face detection domain, a deep investigation into the design of the feature is done mainly on channel types, window size, subsampling method and feature scale. Results of comparative experiments are shown in Figure 6.

Channel types: Three types of channels are used, which are color channel (Gray-scale, RGB, HSV and LUV), gradient magnitude, and gradient histograms. The computation of the latter two channel types could be seen as a generalized version of HoG features. Specifically, gradient magnitude is the biggest response on all three color channels, and oriented gradient histograms follow the idea of HoG in that: 1) rectangular cell size in HoG equals the subsampling factor in aggregated channel features; 2) each orientation bin results in one feature channel (6 orientation bins are used in this paper). Figure 6 (a)~(c) show how much each of these three types alone contributes to the performance of face detection. It can be seen that the gradient histograms contribute most to the performance among all three channel

¹<http://testsetlrs.icg.tugraz.at/research/aflw/>

²<http://www.ics.uci.edu/~x Zhu/face/>

³The Jaccard index is defined as the size of the intersection divided by the size of the union of the sample sets.

types. Figure 6 (d) shows the performances of combinations of these three types computed on different color channels.

Detection window size: Detection window size is the scale to which we resize all face and non-face samples and then train our detector. Larger window size includes more pixels in feature pool and thus may improve the face detection performance. On the other hand, too large window will miss some small faces and diminish the detection efficiency. Figure 6 (e) shows comparison of window size ranging from 32 to 112 with a stride of 16 pixels.

Subsampling: The factor for subsampling can be regarded as the perceptive scale for that it controls the scale at which the aggregation is done. Changing the factor from large to small leads to the feature representation shifting from coarse to fine and the feature pool size getting bigger. Experiments on different subsampling factors are shown in Figure 6 (f). In original aggregate channel features, the way to do subsampling is average pooling. Following the idea in Convolutional Neural Networks, another two ways of subsampling, max pooling and stochastic pooling [24] are tested in Figure 6 (g).

Smoothing: As described in feature description, both pre and post smoothing is done in default setting of aggregate channel features. A binomial filter with a radius of 1 is used for smoothing. The smoothing procedure also has a great influence on the scale of the feature representation. Concretely, pre-smoothing determines how far the local neighborhood is in which local correlations are encoded before channel computation, while post-smoothing determines the neighborhood size in which the computed channel features are integrated with each other. In [6], the former corresponds to the ‘local scale’ of the feature, while the latter represents the ‘integration scale’. We vary the filter radius used in pre and post smoothing and find that both using a radius of 1 gets the best results. Figure 6 (h)~(i) present the comparative results.

Multi-scale: In aggregate channel features, although hidden information at different scale could be extracted at a cost of more weak classifiers, it would be better to make the integrated channel features multi-scaled and thus make themselves more discriminant. Therefore the same or better classification performance can be achieved with fewer weak classifiers. In this part, we implement three multi-scale version of aggregate channel features in the aforementioned three kinds of scale, perceptive scale (subsampling), local scale (pre-smoothing) and integration scale (post-smoothing) and compare their performances. See results in Figure 6 (j)~(l).

Summary: The color channel, gradient magnitude and gradient histograms prove themselves a good match in aggregate channel features. However, different choices of color channel used and on which gradients are computed have a great impact on performance. According to the experi-

ments, LUV channel and gradient magnitude and 6-bin histograms computed on RGB color space (in total 10 channels) are the best choice for face detection.

Larger detection window size generally gets better performance, but will miss many small faces in testing and lead to inefficient detection. In this work, we set the size to 80×80 as its optimal performance.

A subsampling factor of 4 is most reasonable according to the experiments, while different pooling methods show small differences. However, max pooling and stochastic pooling are much slower than average pooling, therefore the average pooling becomes the best match for the sake of efficiency. In this way, the resulting feature pool size of our face detector is $(80/4) \times (80/4) \times 10 = 4000$, considerably smaller than that in VJ framework.

As for multi-scale version of aggregate channel features, multi-local-scale with an additional scale of radius 2 shows the best performance. The probable reason is that pre-smoothing controls the local scale of the neighborhood feature correlations and therefore matches the intuition inside multi-scale best. Compared with other fine-tuning, the multi-scale version has a notable performance gain for that it makes up for the scale uniformity caused by subsampling to some extent. One main drawback is that it doubles the feature pool size and as a result slows down the detection speed somewhat. Based on the trade-off, we implement two face detectors with different scale settings, one is single-scaled with faster speed and the other is multi-scaled with better accuracy. We evaluate and discuss the performances of these two versions in detail in section 5.

3.4. Training design

Besides careful design of the aggregate channel features, experiments on the training process which is similar to that in VJ framework are also carried out. The differences are that the weak classifier is changed into depth-2 decision tree and soft-cascade [1] structure is used. Details of the training design are as follows.

Number of weak classifiers: Given a feature pool size of 4,000, we vary the number of weak classifiers contained in the soft-cascade. In Figure 3 performances of various numbers of weak classifiers ranging from 32 to 8192 are displayed, which shows that apparently more classifiers generate better performance, and when the number gets larger the performance begins to saturate. Since more classifiers slow down the detection speed, there’s a trade-off between accuracy and speed. Searching for the saturate point as the optimal is significant during training in such framework.

Training data: Empirically, more training data will get better performance given powerful representation capacity. In this case, AFLW database is used as the only positive training data. However, as images in AFLW database are

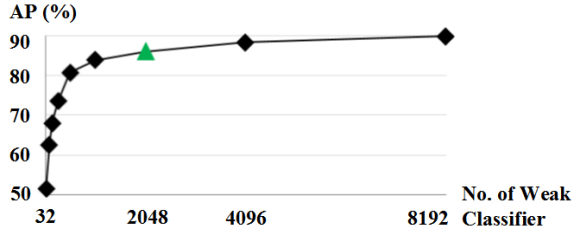


Figure 3. Comparison of different numbers of weak classifier in the soft cascade.

very salient and the background has very less variance, negative samples cropped from the AFLW database can't represent the real world scenario well, which limits the face detection performance in the wild. In this part, we further use PASCAL VOC 2007 database and randomly crop windows from images without person as the new negative samples. Experiments show that the new training data containing cluttered background significantly improve the performance with 4.1%.

Summary: Based on observations above, we choose 2048 as the number of weak classifiers contained in the soft cascade. As each weak classifier is a depth-2 decision tree, it takes only two comparing operations to apply a weak classifier, which is quite fast. During training, as negative data is large, we adopt a standard Bootstrap procedure to sample hard negative samples from PASCAL VOC 2007 in the implementation of the proposed face detector.

4. Multi-view detection

Human faces in real world usually have highly varied poses. In AFLW database, the human pose is divided into three aspects: 1 in-plane rotation 'roll' and 2 out-of-plane rotations 'yaw' and 'pitch'. Because of this large variance in face pose, it is difficult to train a single view face detector to handle all the poses effectively. A multi-view detection is further examined in this part. Due to the adoption of soft-cascade structure, a multi-view version of face detector won't cause too much computation burden. Typically, we divide the out-of-plane rotation yaw into different views and let the classifier itself tolerate the pose variance in the other two types of rotations.

Adopting multi-view detection also brings about many troublesome issues. If handled improperly, the performance will differ greatly. First, detectors of different view will each produce a set of candidate positive windows followed with a set of confidence scores. For application purpose, we need to merge these detections from different views and also remove duplicated windows. A typical approach is Non-Maximum Suppression (NMS) [3]. An issue rises on how to compare confidence scores from different classifiers and how to do window merging in the trade-off between high

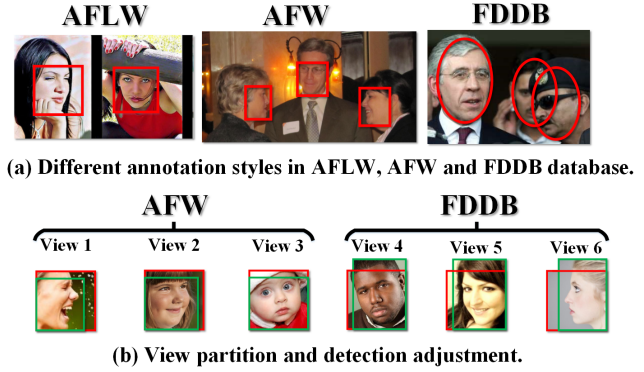


Figure 4. Illustration of different ground-truth annotation styles in databases, the view partition and symmetric detection adjustment. Rectangles with red and green color correspond to detections before and after adjustment.

precision rate and high detection rate. Second, as for detection evaluation, usually the overlap of bounding boxes is used as the criterion. However, annotations in different data sets may not have a consistent style (Figure 4 (a)). This diversity suffers more in profile faces. Since our face detector is trained and tested on different data sets, this issue impacts the performance a lot. Third, detectors of different views need to be trained with different samples separately. How to divide the views therefore becomes another concerning problem. In this section, we address the above three issues successfully by careful designs and therefore fully exploit the advantage of multi-view detection.

4.1. View partition

In the scenario of detecting faces in the wild, pose variation caused by yaw is usually severer than pitch and roll. Therefore we divide the faces in AFLW database according to yaw angle. We have 6 subviews which are horizontally symmetric (see Figure. 4 (b)) because we flip each image in the training set. Specifically, there are 6630, 8446, 9610, 9610, 8446, 6630 images in views from 1 to 6. Benefitting from the symmetry of our model, we can only train three subview detectors of the right side for simplicity, and use these trained right-side detectors to generate the left-side detectors. Detections of all six detectors are then merged to get the final detections. Though multi-view detection significantly improves the detection performance (especially the recall rate), the post-processing of detections from different detectors becomes a trouble. If handled improperly, the performance degrades a lot.

4.2. Post-processing

Difficulties in the post-processing of multi-view detection mainly reflect on the following aspects: 1) different score distributions and; 2) different bounding box styles.

Concretely, as each subview detector is trained separately, their output confidence scores usually have different distributions. What’s more, due to the annotation rule in the AFLW database that the face’s nose is approximately at the center location of the bounding box ground-truth, as the subview changes, the bounding box shifts. This bounding box offset causes difficulty both in detection merging and final evaluation using Jaccard index metric. To solve these annoying issues and make the best use of multi-view detection, we introduce the following methods for post-processing.

Score re-ranking: We propose the following three kinds of score re-ranking: 1) normalizing scores of different views to [0, 1]; 2) defining a new score that has uniform distribution and; 3) taking overlapping detections into consideration.

Normalization: After training a classifier, calculate the output range of the classifier and use the range to do normalization later so that output score has a range of [0, 1].

NewScore: Originally, each weak classifier in the soft-cascade owns a score and final score is the sum of all scores. Instead, we use the number of weak classifier that the image patch passed positively as the new score. Therefore the upper limit of the new score is 2048 in our case.

OverlapRerank: Given an image, multiple detections from multi-view detectors exist each with a score. For each detection, we first calculate the number of overlapped detection it has (overlap threshold is 0.65) and then multiply score of each detection with a factor of its overlapping number ranking¹.

SumofOverlap: Instead of using overlapping as a multiply factor, here we use the sum of overlapped detections’ scores as the current detection’s new score.

Detection merging: Apart from the *Greedy** version of Non Maximum Suppression [6], we also use the detection combination introduced in [20]. It averages the locations of overlapped detections rather than suppresses them.

Detection adjustment: As shown in Figure 4 (a), different databases have different annotation styles of ground-truth. Specifically, AFLW has square annotations with nose located approximately at the center. AFW uses tight rectangular bounding boxes as annotations with the eye-brow being the approximate upper bound. FDDB uses elliptical annotations bounding the whole head. As our detector is trained on AFLW and tested on AFW and FDDB, there exist offsets in both detection position and scale. According to observations, the offsets vary as face pose changes. Therefore we adopt a view-specific detection adjustment to alleviate the offsets. Note that the adjustment is constant for all images and faces in the same database, see Figure 4 (b) for

¹A toy example: Det1: score: 10, nOverlap: 10; Det2: score: 9, nOverlap: 20; Det3: score: 5, nOverlap: 5. After score re-ranking: Det1: score: $10 \times \frac{2}{3} = 6.67$; Det2: score: $9 \times \frac{3}{3} = 9$; Det3: score: $5 \times \frac{1}{3} = 1.67$.

Reranking	AP (%)	Merging	AP (%)
None	91.7	Greedy* NMS	91.7
Normalization	93.5	Combination	93.4
NewScore	92.9		
OverlapRerank	95.0		
SumofOverlap	93.7		

Table 1. Comparison of different methods of score re-ranking and detection merging.

details.

Summary: According to experimental results (Table 1), *OverlapRerank* seems to be the best score re-ranking method. The underlying reason may be that true positives usually have many overlapped detections, while the false positives would only get a few responses. Therefore leveraging this overlapping information in score re-ranking can reduce many false positives. However, in practice, overlap related methods and detection combination both cost much time to process, which is infeasible in a large majority of applications. We finally adopt *Normalization* score re-ranking combined with *Greedy** Non Maximum Suppression for the sake of detection speed.

5. Experiments

In this section, we compare our method with state-of-the-art methods on AFW and FDDB databases which contain challenging faces in the wild. In AFW, we compare with three commercial systems (Google Picasa, Face.com and Face++) and five academic methods (Shen *et al.* [21], Zhu *et al.* [27], DPM [7], multiHOG [27] and Kalal *et al.* [9]). In FDDB, we compare with one commercial system (Olaworks) and six academic methods (Yan *et al.* [23], Boosted Exemplar *et al.* [12], SURF multi-view [13], PEP-Adapt [11], XZJY [21] and Zhu *et al.* [27]) listed on FDDB results page¹.

5.1. Evaluation on benchmark face database

As shown in Figure 5, in AFW, our multi-scale detector achieves an ap value of 96.8%, outperforming other academic methods by a large margin. When it comes to commercial systems, ours is better than Face.com and almost equal to Face++ and Google Picasa. Note that most of our false positives on AFW database are faces that haven’t been annotated (small, seriously occluded or artificial faces like mask and cartoon character).

When evaluated on FDDB database, we follow the evaluation protocol in [8] and report the average discrete and continuous ROC of the ten subfolders. For equality, we fix the number of false positives to 284 (equivalent to an average of 1 False Positive Per Image) and compare the

¹<http://vis-www.cs.umass.edu/fddb/results.html>

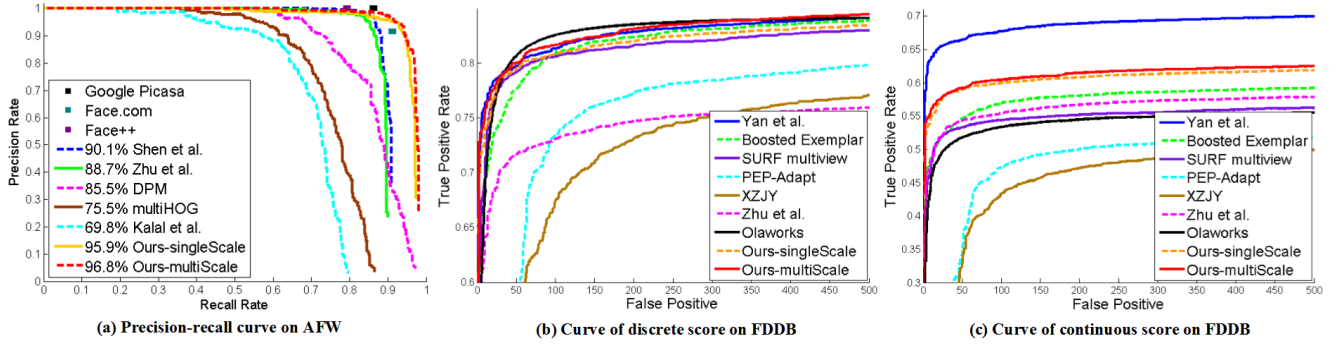


Figure 5. Experimental results on AFW and FDDB database. Best viewed in color.

true positive rate. In discrete score where evaluation metric is the same as in AFW, our detector achieves 83.7%, which is a little better than Yan *et al.* [23]. Note that the ground-truth in FDDB are elliptical faces, therefore the evaluation metric of an overlap ratio bigger than 0.5 cannot reveal the true performance of the proposed detector well. When using continuous score which takes the overlap ratio as the score, our method gets 61.9% true positive rate at 1 FPPI for multi-scale version, surpassing other methods which output rectangular detections by a notable margin (the Yan *et al.* detector outputs the same elliptical detections as the ground-truth, therefore having advantages with this metric). Our detector using single-scale features performs a little worse with the benefit of faster detection speed.

5.2. Discussion

Training efficiency: We implement the method with Piotr’s MATLAB toolbox [4] on a PC with Intel Core i7-3770 CPU and 16GB RAM. With 21,328 positive images and 5,771 negative images in total 6 views, the training process takes about 5.3 mins for a single-scale subview detector containing 2048 weak classifiers and 10.2 mins for multi-scale version. Note that we use much fewer training data than SURF multiview [13] whilst still outperforming their performance.

Comparative results: When inspecting detections of the proposed face detector and other algorithms on the test-sets, some patterns can be found to explain why our detector outperforms others. One evident strength lies in detecting faces with extreme poses. Because we adopt multi-view detection and train each subview detector separately, our detector handles pose variations very well. Second is the outstanding illumination invariance of our detector, which is mainly owing to the extension of channel types to LUV color space and gradient-related channels.

Detection speed: Due to the simple form of aggregate channel features and fast computation of feature pyramid [5], detection is quite efficient. For full yaw pose face detection in VGA image, the proposed detector using

single-scale features runs at 20 FPS on a single thread and 62 FPS if 4 threads are used. If only frontal faces are concerned, the detector runs at 34 FPS and 95 FPS after parallelization. When it comes to the proposed detector using multi-scale features, the above four indices reduce to 15, 42, 21 and 55 FPS. Considering the large performance gain and similar speed, the proposed method can replace Viola-Jones detector for face detection in the wild.

6. Conclusion

A novel feature representation called aggregate channel features possesses the merits of fast feature extraction and powerful representation capacity. In this paper, we successfully apply the feature representation to face detection domain through a deep investigation into the feature design, and propose a multi-scale version of feature which further enriches the representation capacity. Combined with our efforts into solving issues concerning multi-view detection, the proposed multi-view face detector shows state-of-the-art performance in both effectiveness and efficiency on faces in the wild. The proposed method appeals to real world application demands and has the potential to be embedded into low power devices.

Acknowledgement

This work was supported by the Chinese National Natural Science Foundation Projects #61105023, #61103156, #61105037, #61203267, #61375037, National Science and Technology Support Program Project #2013BAK02B01, Chinese Academy of Sciences Project No. KGZD-EW-102-2, and AuthenMetric R&D Funds.

References

- [1] L. Bourdev and J. Brandt. Robust object detection via soft cascade. In *CVPR*, 2005. 1, 3, 4
- [2] S. C. Brubaker, J. Wu, J. Sun, M. D. Mullin, and J. M. Rehg. On the design of cascades of boosted ensembles for face detection. *IJCV*, 2008. 2

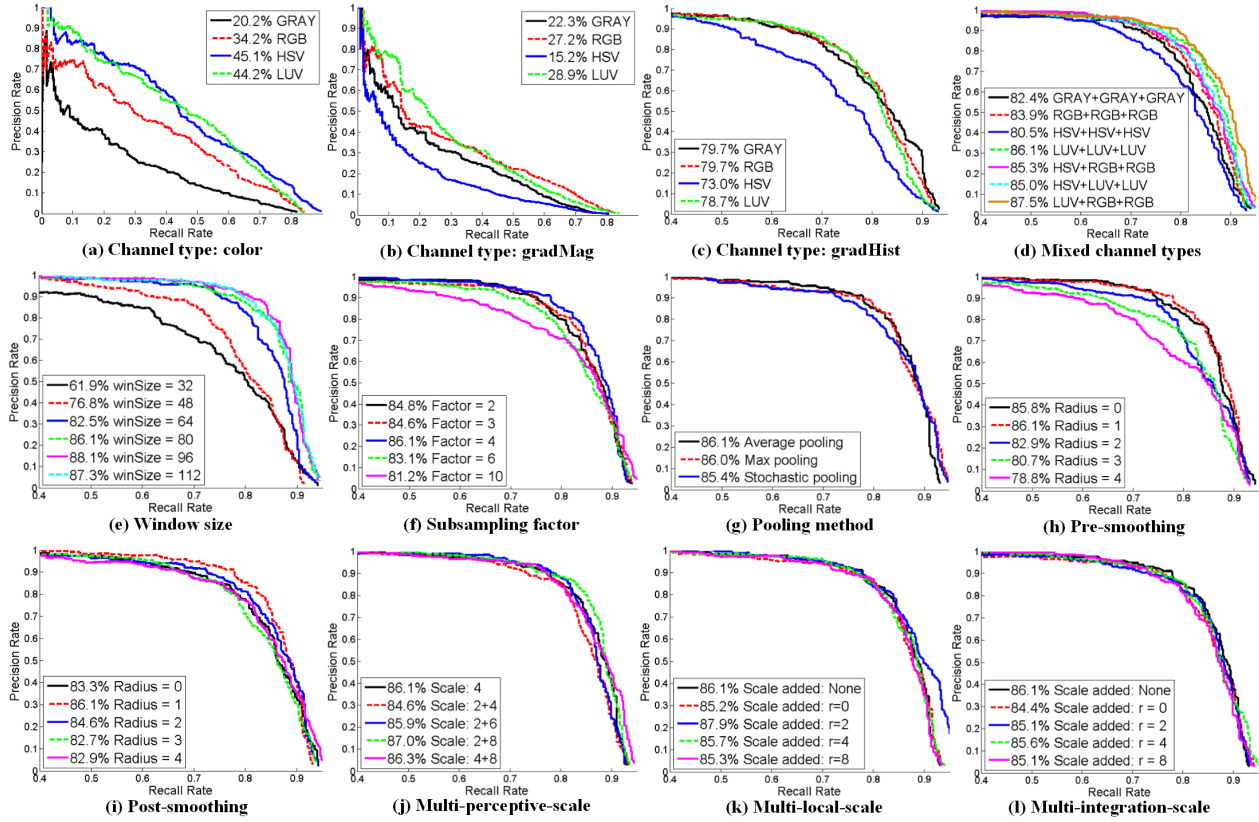


Figure 6. Results of comparative experiments in feature design. Best viewed in color.

- [3] N. Dalal and B. Triggs. Histograms of oriented gradients for human detection. In *CVPR*, 2005. 5
- [4] P. Dollár. Piotr's image and video matlab toolbox (pmt). *Software available at: <http://vision.ucsd.edu/~pdollar/toolbox/doc/index.html>*, 2012. 7
- [5] P. Dollár, R. Appel, S. Belongie, and P. Perona. Fast feature pyramids for object detection. *TPAMI*, 2014. 1, 2, 7
- [6] P. Dollár, Z. Tu, P. Perona, and S. Belongie. Integral channel features. In *BMVC*, 2009. 1, 2, 3, 4, 6
- [7] P. F. Felzenszwalb, R. B. Girshick, D. McAllester, and D. Ramanan. Object detection with discriminatively trained part-based models. *TPAMI*, 2010. 6
- [8] V. Jain and E. Learned-Miller. Fddb: A benchmark for face detection in unconstrained settings. Technical report, 2010. 6
- [9] Z. Kalal, J. Matas, and K. Mikolajczyk. Weighted sampling for large-scale boosting. In *BMVC*, 2008. 6
- [10] M. Koestinger, P. Wohlhart, P. M. Roth, and H. Bischof. Annotated facial landmarks in the wild: A large-scale, real-world database for facial landmark localization. In *First IEEE International Workshop on Benchmarking Facial Image Analysis Technologies*, 2011. 3
- [11] H. Li, G. Hua, Z. Lin, J. Brandt, and J. Yang. Probabilistic elastic part model for unsupervised face detector adaptation. In *ICCV*, 2013. 6
- [12] H. Li, Z. Lin, J. Brandt, X. Shen, and G. Hua. Efficient boosted exemplar-based face detection. In *CVPR*, 2014. 6
- [13] J. Li and Y. Zhang. Learning surf cascade for fast and accurate object detection. In *CVPR*, 2013. 2, 6, 7
- [14] S. Z. Li, L. Zhu, Z. Zhang, A. Blake, H. Zhang, and H. Shum. Statistical learning of multi-view face detection. In *ECCV*, 2002. 1
- [15] S. Liao, A. K. Jain, and S. Z. Li. Unconstrained face detection. Technical report, 2012. 3
- [16] M. Mathias, R. Benenson, M. Pedersoli, and L. Van Gool. Face detection without bells and whistles. In *ECCV*, 2014. 2
- [17] T. Mita, T. Kaneko, and O. Hori. Joint haar-like features for face detection. In *ICCV*, 2005. 1
- [18] M.-T. Pham and T.-J. Cham. Fast training and selection of haar features using statistics in boosting-based face detection. In *ICCV*, 2007. 2
- [19] M.-T. Pham, Y. Gao, V.-D. Hoang, and T.-J. Cham. Fast polygonal integration and its application in extending haar-like features to improve object detection. In *CVPR*, 2010. 1
- [20] P. Sermanet, D. Eigen, X. Zhang, M. Mathieu, R. Fergus, and Y. LeCun. Overfeat: Integrated recognition, localization and detection using convolutional networks. *arXiv:1312.6229v4*. 6
- [21] X. Shen, Z. Lin, J. Brandt, and Y. Wu. Detecting and aligning faces by image retrieval. In *CVPR*, 2013. 6
- [22] P. Viola and M. J. Jones. Robust real-time face detection. *IJCV*, 2004. 1, 2
- [23] J. Yan, Z. Lei, L. Wen, and S. Z. Li. The fastest deformable part model for object detection. In *CVPR*, 2014. 6, 7
- [24] M. D. Zeiler and R. Fergus. Stochastic pooling for regularization of deep convolutional neural networks. *arXiv:1301.3557v1*. 4
- [25] C. Zhang and P. A. Viola. Multiple-instance pruning for learning efficient cascade detectors. In *NIPS*, 2007. 1
- [26] Q. Zhu, M.-C. Yeh, K.-T. Cheng, and S. Avidan. Fast human detection using a cascade of histograms of oriented gradients. In *CVPR*, 2006. 1, 2
- [27] X. Zhu and D. Ramanan. Face detection, pose estimation, and landmark localization in the wild. In *CVPR*, 2012. 6