

Robust Feature Selection Based Lossless Video Compression OF Tiny Video Scenes Using Multi Feature Reduction Technique and Wavelet Transform

S.Ponlatha^{#1} and Dr. R.S. Sabeenian^{*2}

1Associate Professor/ ECE,AVS Engineering College, Ammapet, Salem-636003.

2Professor/ECE & Centre- Head,SIPRO, Sona College of Technology,Salem

Abstract:

The process of video compression has been studied well in recent years, but most approaches are suffers with the precision of compression and decompression quality in reproducing the original video. We propose a novel feature selection approach for video compression which uses feature reduction techniques and wavelet transform. The proposed method identifies scenes and splits them into tiny video scenes. From identified scenes, the proposed method extracts the object maps and selects the features which are varying at subsequent scenes. For each scene from the video, we generate set of feature points which represent the foreground feature of the scene. The object map has all the features present in the video and used for restoring the original video. The feature set is reduced by identifying the common background and foreground objects and their features. Only the varying snapshot are transferred and other snapshot are restored with the same set of features. The proposed method has produced efficient compression ratio than earlier approaches.

Index Terms: Video Compression, Robust Features, Wavelet Transform.

Introduction:

The growth of internet technology tends people to share and send variety of information between the user. The transfer of video clip need more bandwidth constraint and increases the latency also. Now a day video conferencing, a popular application used in real time environment which need the clippings without any delay. Sending such video clips needs more sophisticated methodology to provide better

service. So that, the video has to be compressed and relayed to reduce the bandwidth consumption and reducing the latency.

The video compression is the process of reducing the special property and keeping the features as its and originality. There are many video compression strategies available and proposed, for example video compression using wavelet compression methods. The wavelet transform compress the original signals into low level signals and compress the signals with originality which is popular in this problem. The problem of accuracy is the question here where the wavelet transform produces quality output but the time complexity is higher.

The speed up robust features which is a digital processing method published, to support image retrieval methods. In that approach, the input image has been converted into number of sub sampled image and identifies set of interest points where there exist more features around a point. Identified points are converted into feature descriptors using hessian matrix, around 64 features are extracted to generate feature descriptor. This approach supports varying shape and other features which could be applicable for various image processing tasks. We trust the approach and could adapt the method for our problem with little changes in this paper.

Related Works:

There are number of video compression method has been discussed for efficient compression of videos and we discuss most of them here around the problem statement.

A Novel Approach towards Video Compression for Mobile Internet using Transform Domain Technique [3], proposed to reduce the data rates required to transfer multimedia files on the mobile networks, we require one compression algorithm which gives us higher compression with easiness of implementation. So here the paper present one algorithm which is much simpler to implement, as it is based on MPEG-2 then also gives compression in considerable amount, which is comparable to MPEG-4 techniques.

Multiframe adaptive Wiener filter super-resolution with JPEG2000-compressed images [4], propose and compare novel processing architectures for applying multiframe SR with JPEG2000 compression. We propose a modified adaptive Wiener filter (AWF) SR method and study its performance as JPEG2000 is incorporated in different ways. In particular, we perform compression prior to SR and compare this to compression after SR. We also compare both independent-frame compression and difference-frame compression approaches. We find that some of the SR artifacts that result from compression can be reduced by decreasing the assumed global signal-to-noise ratio (SNR) for the AWF SR method. We also propose a novel spatially adaptive SNR estimate for the AWF designed to compensate for the spatially varying compression artifacts in the input frames.

Yuming Fang, A Video Saliency Detection Model in Compressed Domain [10], propose a novel video saliency detection model based on feature contrast in compressed domain. Four types of features including luminance, color, texture, and motion are extracted from the discrete cosine transform coefficients and motion

vectors in video bitstream. The static saliency map of unpredicted frames (I frames) is calculated on the basis of luminance, color, and texture features, while the motion saliency map of predicted frames (P and B frames) is computed by motion feature. A new fusion method is designed to combine the static saliency and motion saliency maps to get the final saliency map for each video frame. Due to the directly derived features in compressed domain, the proposed model can predict the salient regions efficiently for video frames.

Scanned Document Compression Using Block-Based Hybrid Video Codec [11], proposes a hybrid pattern matching/transform-based compression method for scanned documents. The idea is to use regular video interframe prediction as a pattern matching algorithm that can be applied to document coding. We show that this interpretation may generate residual data that can be efficiently compressed by a transform-based encoder. The efficiency of this approach is demonstrated using H.264/advanced video coding (AVC) as a high-quality single and multipage document compressor. The proposed method, called advanced document coding (ADC), uses segments of the originally independent scanned pages of a document to create a video sequence, which is then encoded through regular H.264/AVC.

Performing scalable lossy compression on pixel encrypted images [13], propose a scalable lossy compression scheme for images having their pixel value encrypted with a standard stream cipher. The encrypted data are simply compressed by transmitting a uniformly subsampled portion of the encrypted data and some bitplanes of another uniformly subsampled portion of the encrypted data. At the receiver side, a decoder performs content-adaptive interpolation based on the decrypted partial information, where the received bit plane information serves as the side information that reflects the image edge information, making the image reconstruction more precise. When more bit planes are transmitted, higher quality of the decompressed image can be achieved.

Royalty free video coding standards in MPEG [15], report on the recent developments in royalty-free codec standardization in MPEG, particularly Internet video coding (IVC), Web video coding (WVC), and video coding for browser, by reviewing the history of royalty-free standards in MPEG and the relationship between standards and patents. A Video Saliency Detection Model in Compressed Domain [17], propose a novel video saliency detection model based on feature contrast in compressed domain. Four types of features including luminance, color, texture, and motion are extracted from the discrete cosine transform coefficients and motion vectors in video bitstream. The static saliency map of unpredicted frames (I frames) is calculated on the basis of luminance, color, and texture features, while the motion saliency map of predicted frames (P and B frames) is computed by motion feature. A new fusion method is designed to combine the static saliency and motion saliency maps to get the final saliency map for each video frame. Due to the directly derived features in compressed domain, the proposed model can predict the salient regions efficiently for video frames.

In [17], the author proposes two schemes that balance energy consumption among nodes in a cluster on a wireless video sensor network. In these schemes, we divide the compression process into several small processing components, which are

then distributed to multiple nodes along a path from a source node to a cluster head in a cluster. We conduct extensive computational simulations to examine the truth of our method and find that the proposed schemes not only balance energy consumption of sensor nodes by sharing of the processing tasks but also improve the quality of decoding video by using edges of objects in the frames.

All the above discussed approaches has the problem of feature restoration and compression efficiency. To overcome these issues , we propose a naval approach which uses wavelet transform and object graphs for video compression.

Proposed Approach:

The proposed naval approach has various stages of compression and they are Tiny Video Scene Generation, Wavelet Compression and Robust Feature Reduction. The tiny video scene generation identifies the scenes from video clip and separate them into number of tiny video scenes which will be used for further processing. At the Wavelet compression , the scenes are converted into snapshots and further compressed using wavelet transform and finally at the feature reduction stage the variant features are identified and used for compression.

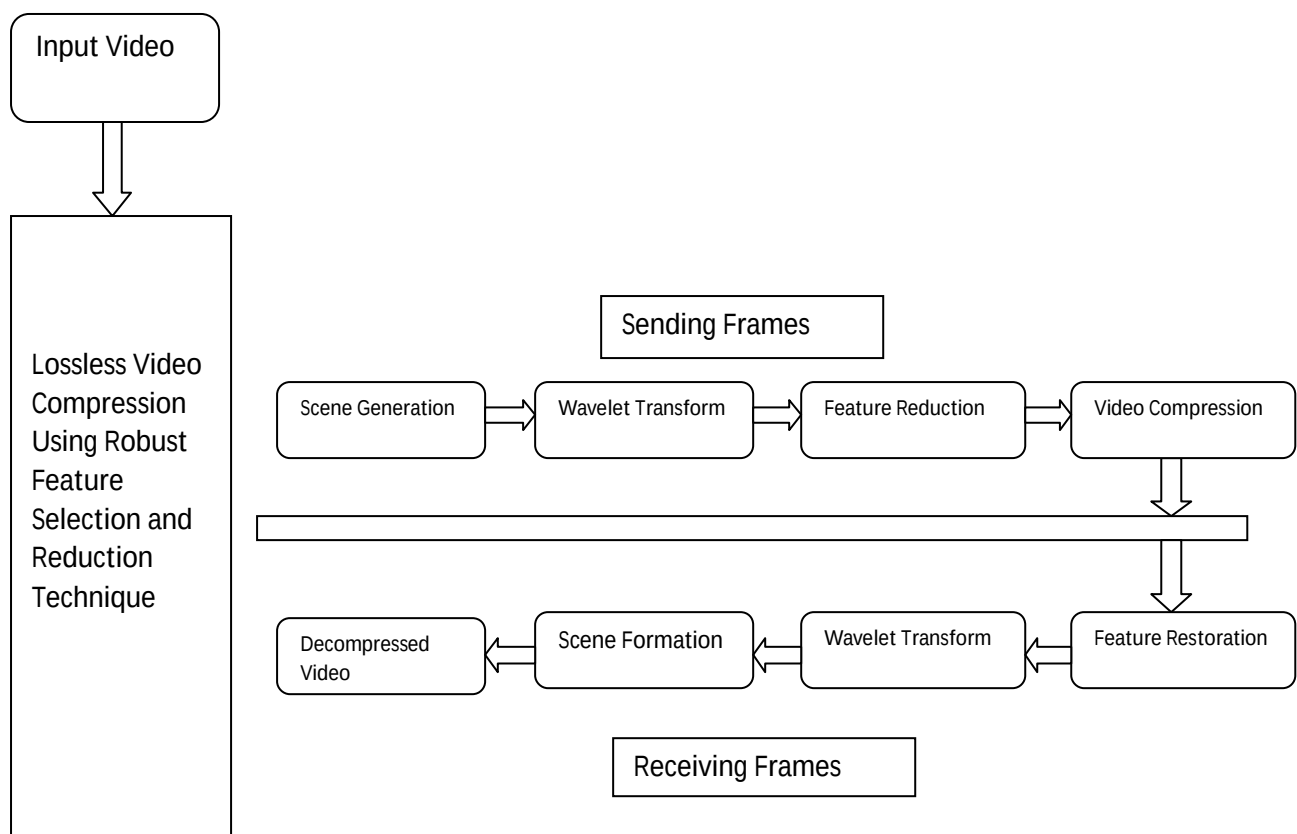


Figure1: Proposed System Architecture.

Tiny Video Scene Generation:

The input video clip is converted into number of snapshots according to the quality of the video. Generally we generate 40 snapshots per second of the video time. For each subsequent image, we identify SURF based interest points. Unlike popular SURF, we simply generate the interest points with the same size of image window to reduce the time complexity. From identified interest points, the hessian matrix is computed to generate the feature descriptors. We compute the similarity measure (Euclidean Distance) between consequent set of feature descriptors to identify the level of deviation present between snapshots of the video clip. If the deviation between two consequent frames becomes more than a limit then we consider that the frame starts a new video. Otherwise, the method considers the snapshot as in the same video scene and keeps the feature descriptors for feature usage.

Algorithm:

Input: Video V_i .

Output: Scene set S .

step1: start

step2: initialize scenery set S , Frame set F_s

step3: for each second of video V_i

generate snapshot S_s . //Here we generate snapshots at particular second

$FS = FS + \sum (\int_1^N (V \times M))$. // F_s is sum of all snapshot of video at all

seconds

M - number of frames per second.

end.

step4: for each frame F_i from F_s .

Compute interest point I_p for F_i .

$IP = \int \sum (I_p(F_i))$

Feature Descriptor $FD = \int_{i=1}^{size(IP)} \sum Feature(I_p(i))$

Compute Euclidean distance with the previous frames IP descriptor.

$Ed = \int_{i=1}^{size(FD)} Dist(FD(i) - FD(i - 1))$

If $ED > FD\text{-Threshold}$ Then

keep original

else

Perform background subtraction.

$F_i = \int F_i \cap (F_i \times F_i.Alpha)$

end

step5: stop.

The tiny scene generation algorithm generates number of frames for each second of the video and with the each snapshot we identify the feature descriptors and compute the Euclidean distance between the feature descriptors of the previous frames. If the distance is greater than the previous frame then it is considered as the first frame of the scene and a new scene is initialized. Otherwise the frame is

considered as the part of the existing scene and background subtraction is performed and added to the existing scene.

Wavelet Compression:

The wavelet transform is performed at multi level over identified separate tiny video scenes. Each video scene is converted into number of snapshots and each still image is applied with wavelet compression. The wavelet transformed image is used for further processing. The wavelet transform is performed in multilevel which retains the original features of the image and could be restored at later stage. The precision of wavelet transform is well known and used for our methodology.

Robust Feature Reduction:

The wavelet transformed image is applied with background subtraction and the foreground object is used to generate object map. The object map consists of various features and neighbors connected according to them. We compute the set of objects maps between subsequent frames of the scene to identify the difference between them. Generated object maps and variance are used to transform the image to the other side of the communication.

Algorithm:

Input: Video Scenes and Frames set Fs.

Output: Feature Reduced maps.

Step1: for each scene Si from scene set SS

For each frame Fi From Fs

Generate Object Maps $Omp = \int Fi(\sum_{i=1}^N Fi.pi \cup Fi.pj)$

if Fi == Starting Frame

else

Compute difference with previous map.

$$MD = \int_{i=1}^N Dist(Fi - Fj)$$

if MD > Map Threshold Then

Compute Covariant Matrix CM.

$$CM = \int_{i,j=1}^N \forall (Fi(Pi) \cap Fj(Pj))$$

end

end

end

Step2: stop.

The method generates object maps for each frame and from the object map the distance between the previous frames map is computed. Based on the distance value a covariant matrix is computed based on the map reduce threshold. The reduced features are transmitted to the other end to perform video decompression. The

computed covariance matrix shows the difference between subsequent frames which is computed using object maps. The variant pixels of the image only kept alive and others are neutralized to perform wavelet transform.

Video Compression:

The generated features and object maps and video frames are used to compress the image. The object map different between scenes and frames are used to compress the image. The wavelet transformed image which is the initial frame is sent as it is for a reference and the other frames are removed with background subtraction and the varying features with object map only will be sent to the receiver side. The receiving side has to check with the newly generated object map from the received image and has to compare with the previous map. Based on the difference, it has to store the original features. The image with background shows the start of the next scene and it acts accordingly.

Results and Discussion:

We have proposed a novel video compression approach which uses object maps and wavelet transform. The approach has been implemented in Matlab and evaluated for its accuracy and performance with various other approaches. The proposed approach has produced efficient results in all the factors of accuracy and efficiency than other methods.



Figure1: shows the snapshot generated from the input video clip.

The figure 1, shows the single frame or snapshot generated by the proposed method.



Figure 2: shows the result of compressed image with wavelet transform.

The figure 2 shows the result generated by wavelet transform and the resultant image has retained the quality as well.

Image/Video compression results are presented for various test scenarios with the help of a motion picture obtained from geographic sites which contain a set of 98 frames. The set of frames are tested for, the retraining frames method and the self-adaptive method. The results obtained from the above tests were useful in drawing some conclusions regarding the aforementioned techniques. The compression ratio and peak-signal-to-noise ratios (PSNR) are calculated based on the following formulas for all the test scenarios.

$$\text{Compression Ratio} = \frac{(K \times L \times M \times N)}{T}$$

K- Number of variant pixels

L- Level of wavelet transfer

M- Number of pixels

T-Total number of pixels

PSNR=10× log (1/).

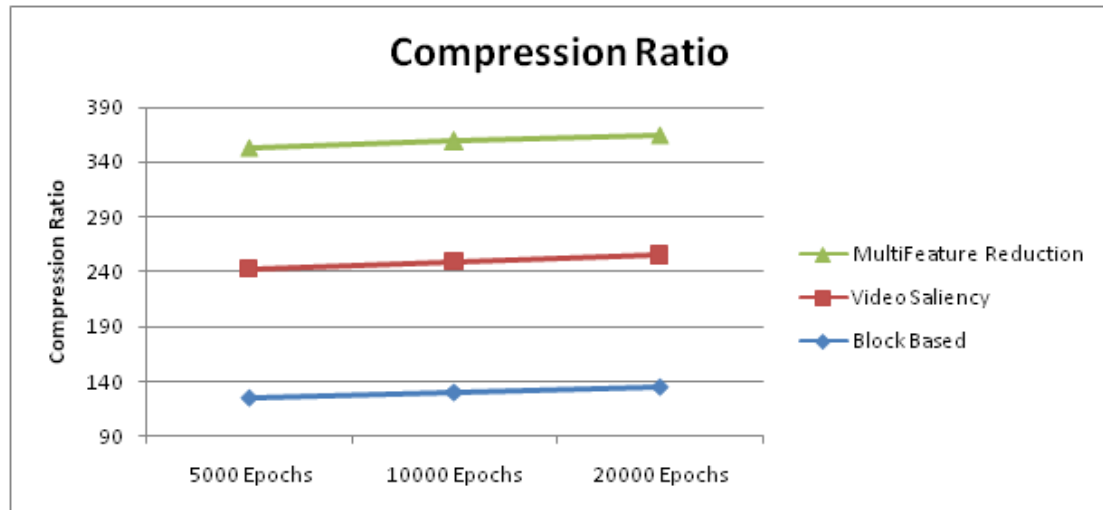
Table1: Comparison of peak signal to noise ratio of different methods.

Number of Epochs	Block Based	Video Saliency	Multi Feature Reduction
5000 Epochs	12.8	15.6	23.2
10000 Epochs	13.9	17.4	28.6
20000 Epochs	15.6	19.2	33.9

The Table 1, shows the comparison of peak signal to noise ration produced by different methods on different epochs. It shows clearly that the proposed Multi Feature Reduction technique has produced higher ratio than other methods.

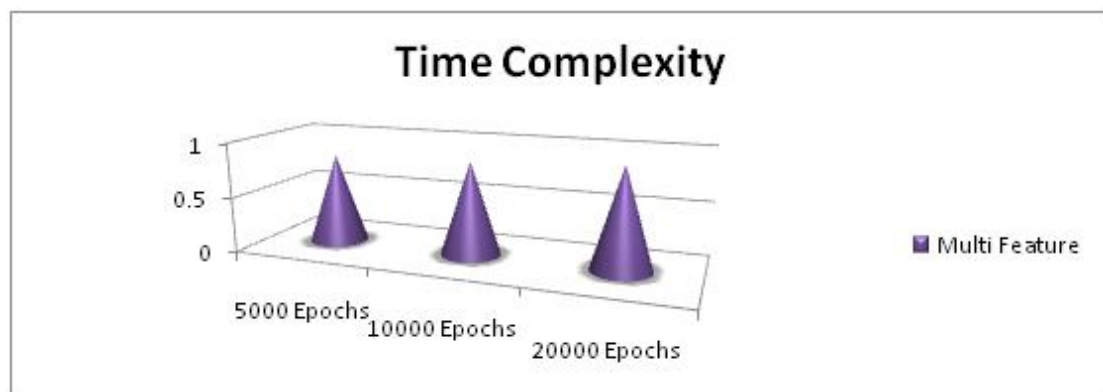
Initially, the “Lion” image is trained for different epochs (100,300,500) and 4 hidden nodes in the network. The network is also tested for a still image with the above parameters. We can see that as the training was increased to 500 epochs, the weights seem to be better adapted to the particular image, and thus the quality of the reconstructed image is higher compared to the one trained for 100 or 300 epochs. Nevertheless, training for 500 epochs has higher processing requirements compared to

the other two cases. Moreover, Graph1 illustrates that as the compression ratio increases; the difference in terms of PSNR between the three different cases becomes negligible.



Graph1: shows the compression ratio achieved.

The graph1, shows the compression ratio generated by different methods at different epochs, for each of the epochs, the frames are generated and compressed. The proposed method has produced efficient compression ratio compared to other methods. The method has been evaluated with more epochs from 1 epoch. In all the cases, the proposed method has produced efficient compression than others.



Graph2: Time complexity graph

The graph2 shows the time taken value of the proposed method for different epochs of jpeg images. It shows that the proposed method has taken less time to

compress the image. The proposed multi-level snapshot algorithm takes less time to compress the video than earlier approaches. To compress a video with 100,300,500 second or epoch , the proposed method takes less time only.

Conclusion:

We proposed a lossless approach for video compression with wavelet transform and the object maps. We split the video into frames and then scene identification is performed using intensity pattern change. Based on identified scenes, we perform wavelet transform which reduces the size of signal. Using the transformed image, we compute the object maps of all the images of a scene and identify the co-variance of subsequent image or frames within the scene. The pixel which varies with intensity is identified and kept original and others are neutralized. This makes the size of frame lesser, because only the variant pixels are used and this reduces the space also. The compressed images are transferred to the other end where they perform the reverse process to restore the same features and video. The proposed approach enhanced the quality of video compression and produces efficient results. The prescribed approach can be further improved by including other map-reduce techniques which can be performed on various other features of the image like shapes. The approach can be further improved by adapting object maps at different layers of a single frame to be transmitted where a single frame has various number of object maps which represent the background changes.

References:

1. Parallelization of Full Search Motion Estimation Algorithm for Parallel and Distributed Platforms, International Journal of Parallel Programming, April 2014, Volume 42, Issue 2, pp 239-264.
2. Yi Gao, Jun Zhou, Motion vector extrapolation for parallel motion estimation on GPU, Multimedia Tools & Applications, 2014
3. Dhaval R Bhojani and Ved Vyas Dwivedi. Article: A Novel Approach towards Video Compression for Mobile Internet using Transform Domain Technique. International Journal of Computer Applications 58(10):29-33, November 2012.
4. Multiframe adaptive Wiener filter super-resolution with JPEG2000-compressed images, EURASIP Journal on Advances in Signal Processing April 2014, 2014:55,
5. Shen M, Xue P, Wang C: Down-sampling based video coding using super-resolution technique. IEEE Trans. Circ. Syst. Video Tech. 2011,21(6):755–765
6. Hardie RC, Barnard KJ: Fast super-resolution using an adaptive Wiener filter with robustness to local motion. Optics. Express 2012,20(19):21053–21073.
7. Hardie RC, Barnard KJ, Raul O: Fast super-resolution with affine motion using an adaptive Wiener filter and its application to airborne imaging. Optics

- Express. 2011,19(27):26208–26231
8. He Q, Schultz R: Super-resolution reconstruction by image fusion and application to surveillance videos captured by small unmanned aircraft systems. In Sensor Fusion and Its Applications. Edited by: Thomas C. Rijkea: InTech; 2010:475–486.
 9. Camrago A, He Q, Palaniappan K, Jara F: Super-resolution mosaics from airborne video using robust gradient regularization. Paper presented in SPIE defense, security and sensing. Burlingame, CA, USA: International Society for Optics and Photonics; 2013
 10. Yuming Fang, A Video Saliency Detection Model in Compressed Domain [10], , IEEE Transaction on Circuit and Systems for video technology, vol:24 , issue:1, pp:27-38, 2014.
 11. Zaghetto, Scanned Document Compression Using Block-Based Hybrid Video Codec, IEEE Transaction on image processing, vol.26, issue.6,pp:2420-2428, 2013.
 12. Ohm J, High efficiency video coding: the next frontier in video compression [Standards in a Nutshell], Ieee Transaction on Signal processing and imaging, vol:30, issue.1, pp:152-158, 2013.
 13. Xiangui Kang, Performing scalable lossy compression on pixel encrypted images [13], EURASIP_Journal on Image and Video Processing, May 2013, 2013:32,
 14. Zhang X, Feng G, Ren Y, Qian Z: Scalable coding of encrypted images. IEEE Trans. Image Process. 2012,21(6):3108–3114
 15. Choi K, Royalty-Free Video Coding Standards in MPEG, IEEE Transaction on Signal Processing, vol:31, issue:1, PP:145-155, 2014.
 16. Paurazad,HEVC: The New Gold Standard for Video Compression: How Does HEVC Compare with H.264/AVC?, IEEE Transaction on consumer electronics magazine, vol:1, issue.3, 2012.
 17. Video Compression Schemes Using Edge Feature on Wireless Video Sensor Networks, Journal of Electrical and Computer Engineering Volume 2012, 2012.

