



ISSN: 2321-2152

IJMECE

*International Journal of modern
electronics and communication engineering*

E-Mail
editor.ijmece@gmail.com
editor@ijmece.com

www.ijmece.com

MOVING OBJECT DETECTION IN COMPLEX SCENE USING SPATIOTEMPORAL STRUCTURED-SPARSE RPCA

1.SHIVA PARVATI,2. S NAVYA SRI ,3. SHAIK BEEBI JHONY ,4. P SOUMYA SRI,5. P KALYANI

ABSTRACT:

Moving object detection is a fundamental step in various computer vision applications. Robust Principal Component Analysis (RPCA) based methods have often been employed for this task. However, the performance of these methods deteriorates in the presence of dynamic background scenes, camera jitter, camouflaged moving objects, and/or variations in illumination. It is because of an underlying assumption that the elements in the sparse component are mutually independent, and thus the spatiotemporal structure of the moving objects is lost. To address this issue, we propose a spatiotemporal structured sparse RPCA algorithm for moving objects detection, where we impose spatial and temporal regularization on the sparse component in the form of graph Laplacians. Each Laplacian corresponds to a multi-feature graph constructed over superpixels in the input matrix. We enforce the sparse component to act as eigenvectors of the spatial and temporal graph Laplacians while minimizing the RPCA objective function. These constraints incorporate a spatiotemporal subspace structure within the sparse component. Thus, we obtain a novel objective function for separating moving objects in the presence of complex backgrounds. The proposed objective function is solved using a linearized alternating direction method of multipliers based batch optimization. Moreover, we also propose an online optimization algorithm for real-time applications. We evaluated both the batch and online solutions using six publicly available datasets that included most of the aforementioned challenges. Our experiments demonstrated the superior performance of the proposed algorithms compared with the current state-of-the-art methods.

INTRODUCTION

Intelligent video surveillance is a new research direction in the field of computer vision. It uses the method of computer vision and detects the movement target in the monitoring scene by automatic analysis the image sequence by the camera recording. And the research on moving target detection and extraction algorithm can be said to be key issues in intelligent video. Its purpose is the detection and extraction of the moving targets from the scene of the video image sequence. Therefore the effective detection of moving targets determines the system performance. Therefore, this article focuses on key technology in the moving targets detection and extraction.

In this paper, firstly, it has a brief introduction of pretreatment of the video images. It reduces the error in the image processing after. Secondly the paper focuses on analysis comparison the two algorithms: the background subtraction and the frame difference. Lastly, this paper selects based on the background subtraction method to improve it and present a moving target detection algorithm based on the background which has dynamic changes. In modern battles, long-distance attacking missile develops to intelligent, high precision and remote controllability. Midcourse guidance uses GPS/INS with terrain matching.

1.ASSISTANT PROFESSOR,2,3,4&5 UG SCHOLAR

DEPARTMENT OF ECE, MALLA REDDY ENGINEERING COLLEGE FOR WOMEN, HYDERABAD

Terminal guidance uses radar, infrared imaging technology or infrared imaging technology with data link. Infrared imaging guidance technology can autosearch, auto-capture, auto-identify target, then can autotrace target because there are many features such as high precision, good anti-interference, good concealment capability and so on and it has been research hotspot in accurate terminal guidance field [1]. At present, the infrared seekers has been the second products whose type products are AAWS-M in America and Triget belongs to German, France and Britain. The information captured by infrared seekers usually is serial image [2]. To treat infrared serial images intelligently is the precondition for accurate terminal guidance, and we can make infrared seekers have better tracing target ability. From martial application, region of interest (ROI) of target in serial images is the region in moving target. So the process of automatic extraction of ROI in infrared serial images is the process of detecting moving target then extraction moving target region. It is a hotspot in computer vision fields that to trace target and to extract ROI from serial images with complex background. The technology used in missile guidance, video controller and traffic manager commonly while it also is an important issue for automatic extraction of ROI. There are two methods for extraction ROI: one is human detected regions of interest (hROI) which is selected according to ROI by human, and another is algorithmically detected regions of interest (aROI) which is selected according to characters of the image [3]. This paper mainly studied the target detection algorithm in static scenes and dynamic scenes, automatic extraction algorithm of ROI and image segmentation issues. The result can improve the efficiency of accurate guidance. In a natural scene, objects of interest often move amidst complicated backgrounds that are themselves in motion e.g. swaying trees, moving water, waves and rain. The visual system of animals is well adapted to recognizing the most important moving object (referred to henceforth as the "target"), in such scenes. In fact, this ability is central to survival, for instance, by aiding in the identification of potential predators or prey while ignoring unimportant motion in the background. Apart from the obvious importance in visual systems of the biological world, target detection is extremely useful for various computer vision applications such as object recognition in video,

activity and gesture recognition, tracking, surveillance and video analysis. For instance, a robot or an autonomous vehicle could benefit from a module to identify objects approaching it amidst possibly moving backgrounds like dust storms, to do effective path planning. However, unsupervised moving target detection, often posed as the related problem of background subtraction, is hard to solve using conventional techniques in computer vision (see (Sheikh & Shah, 2005) for a review). Extracting the foreground object moving in a scene where the background itself is dynamic is so complex that even though background subtraction is a classic problem in computer vision, there has been relatively little progress for these types of scenes. A common assumption underlying many techniques for background subtraction is that the camera capturing the scene is static. (Staufer & Grimson, 1999; Elgammal, Harwood, & Davis, 2000; Wren, Azarbayejani, Darrell, & Pentland, 1997; Monnet, Mittal, Paragios, & Ramesh, 2003; Tavakkoli, Nicolescu, & Bebis, 2006). However, this assumption places severe restrictions on the applicability of such techniques to real-world video clips, that are often shot with hand-held cameras or even on a moving platform in the case of autonomous vehicles. Conventional techniques to address this problem involve explicit camera motion compensation (Jung & Sukhatme, 2004), followed by stationary camera background subtraction techniques. But these methods are cumbersome and require a reliable estimate of the global motion. In extreme cases, when the background itself is highly dynamic, a unique global motion itself may not be possible to estimate. Another disadvantage of most current approaches is that they model the background explicitly and assume that the algorithm will initially be presented with frames containing only the background (Monnet et al., 2003; Staufer & Grimson, 1999; Zivkovic, 2004). The background model is built using this data, and regions or pixels that deviate from this model are considered part of the target or foreground. Hence, these techniques are supervised, and the initial phase could be thought of as *training* the algorithm to learn the background parameters. The need to train such algorithms for each scene separately limits their ability to be deployed for automatic surveillance tasks, where manual re-training of the module to operate in each new scene is not feasible.

- Infrared images can represent space distribution of infrared radiances between the target and its background. The follows are the characters of infrared images [4]:
- Infrared images represent temperature distribution of the object. They are gray images and there are not colors or hatching. So there is lower resolution for human.
- There are higher space correlativity and lower contrast for infrared images because of much physical interference.
- The definition of infrared images is lower than visible images because the space resolution and detection ability of infrared imaging system are not as good as visible CCD array.
- There are many noises in infrared images.
- There is a little changing range in gray values of infrared image. So there are obvious wave crest in histogram of infrared image compared with histogram of visible image. This paper made experiment using Lena image and Infrared tank image.

The bounds between target and its background are very blurry and there are many noises in infrared image ecause there are more details in infrared image captured in complex environment. There are obvious temperature differences between the target and the background in infrared image while they have different gray ranges in the image. So we should study target detection algorithms in various situations firstly if we'll extract ROI of target automatically. In a natural scene, objects of interest often move amidst complicated backgrounds that are themselves in motion e.g. swaying trees, moving water, waves and rain. The visual system of animals is well adapted to recognizing the most important moving object (referred to henceforth as the "target"), in such scenes. In fact, this ability is central to survival, for instance, by aiding in the identification of potential predators or prey while ignoring unimportant motion in the background. Apart from the obvious importance in visual systems of the biological world, target detection is extremely useful for various computer vision applications such as object recognition in video, activity and gesture recognition,

tracking, surveillance and video analysis. For instance, a robot or an autonomous vehicle could benefit from a module to identify objects approaching it amidst possibly moving backgrounds like dust storms, to do ative path planning. However, unsupervised moving target detection, often posed as the related problem of background subtraction, is hard to solve using conventional techniques in computer vision(see (Sheikh & Shah, 2005) for a review). Extracting the foreground object moving in a scene where the background itself is dynamic is so complex that even though background subtraction is a classic problem in computer vision, there has been relatively little progress for these types of scenes. A common assumption underlying many techniques for background subtraction is that the camera capturing the scene is static. (Stau_er Grimson, 1999; Elgammal, Harwood, & Davis, 2000; Wren, Azarbayejani, Darrell, & Pentland,1997; Monnet, Mittal, Paragios, & Ramesh, 2003; Tavakkoli, Nicolescu, & Bebis, 2006). However, this assumption places severe restrictions on the applicability of such techniques to real-world video clips, that are often shot with hand-held cameras or even on a moving platform in the case of autonomous vehicles. Conventional techniques to address this problem involve explicit camera motion compensation (Jung&Sukhatme, 2004), followed by stationary camera background subtraction techniques. But these methods are cumbersome and require a reliable estimate of the global motion. In extreme cases, when the background itself is highly dynamic, a unique global motion itself may not be possible to estimate. Another disadvantage of most current approaches is that they model the background explicitly and assume that the algorithm will initially be presented with frames containing only the background (Monnet et al., 2003; Stau_er & Grimson, 1999; Zivkovic, 2004). The background model is built using this data, and regions or pixels that deviate from this model are considered part of the target or foreground. Hence, these techniques are supervised, and the initial phase could be thought of as *training* the algorithm to learn the background parameters. The need to train such algorithms for each scene separately limits their ability to be deployed for automatic surveillance tasks, where manual re-training of the module to operate in each new scene is not feasible.

LITERATURE REVIE

IN “R. ACHANTA, A. SHAJI, K. SMITH, A. LUCCHI, P. FUA, AND S. SUSSTRUNK, ” “SLIC SUPERPIXELS COMPARED TO STATE-OF-THE-ART SUPERPIXEL METHODS,” IEEE T-PAMI, VOL. 34, NO. 11, PP. 2274–2282, 2012” Computer vision applications have come to rely increasingly on superpixels in recent years, but it is not always clear what constitutes a good superpixel algorithm. In an effort to understand the benefits and drawbacks of existing methods, we empirically compare five state-of-the-art superpixel algorithms for their ability to adhere to image boundaries, speed, memory efficiency, and their impact on segmentation performance. We then introduce a new superpixel algorithm, simple linear iterative clustering (SLIC), which adapts a k-means clustering approach to efficiently generate superpixels. Despite its simplicity, SLIC adheres to boundaries as well as or better than previous methods. At the same time, it is faster and more memory efficient, improves segmentation performance, and is straightforward to extend to supervoxel generation. Superpixel algorithms group pixels into perceptually meaningful atomic regions, which can be used to replace the rigid structure of the pixel grid. They capture image redundancy, provide a convenient primitive from which to compute image features, and greatly reduce the complexity of subsequent image processing tasks. They have become key building blocks of many computer vision algorithms, such as top scoring multiclass object segmentation entries to the PASCAL VOC Challenge [9], [29], [11], depth estimation [30], segmentation [16], body model estimation [22], and object localization [9]. There are many approaches to generate superpixels, each with its own advantages and drawbacks that may be better suited to a particular application. For example, if adherence to image boundaries is of paramount importance, the graph-based method of [8] may be an ideal choice. However, if superpixels are to be used to build a graph, a method that produces a more regular lattice, such as [23], is probably a better choice. While it is difficult to define what constitutes an ideal approach for all applications, we believe the following properties are generally desirable: 1) Superpixels should adhere well to image boundaries. 2) When used to reduce computational complexity as a preprocessing step, superpixels should be fast to compute, memory efficient, and simple to use. 3) When used for

segmentation purposes, superpixels should both increase the speed and improve the quality of the results. We therefore performed an empirical comparison of five state-of-the-art superpixel methods [8], [23], [26], [25], [15], evaluating their speed, ability to adhere to image boundaries and impact on segmentation performance. We also provide a qualitative review of these, and other, superpixel methods. Our conclusion is that no existing method is satisfactory in all regards. To address this, we propose a new superpixel algorithm: simple linear iterative clustering (SLIC), which adapts kmeans clustering to generate superpixels in a manner similar to [30]. While strikingly simple, SLIC is shown to yield stateof-the-art adherence to image boundaries on the Berkeley benchmark [20], and outperforms existing methods when used for segmentation on the PASCAL [7] and MSRC [24] data sets. Furthermore, it is faster and more memory efficient than existing methods. In addition to these quantifiable benefits, SLIC is easy to use, offers flexibility in the compactness and number of the superpixels it generates, is straightforward to extend to higher dimensions, and is freely available



Fig. 1: Images segmented using SLIC into superpixels of size 64, 256, and 1024 pixels (approximately).

IN “H. BHASKAR, L. MIHAYLOVA, AND A. ACHIM, “VIDEO FOREGROUND DETECTION BASED ON SYMMETRIC ALPHA-STABLE MIXTURE MODELS,” IEEE T-CSVT, VOL. 20, NO. 8, PP. 1133–1138, 2010” Background subtraction (BS) is an efficient technique for detecting moving objects in video sequences. A simple BS process involves building a model of the background and extracting regions of the

foreground (moving objects) with the assumptions that the camera remains stationary and there exist no movements in the background. These assumptions restrict the applicability of BS methods to real-time object detection in video. In this letter, we propose an extended cluster BS technique with a mixture of symmetric alpha-stable (SaS) distributions. An online self-adaptive mechanism is presented that allows automated estimation of the model parameters using the log moment method. Results over real video sequences from indoor and outdoor environments, with data from static and moving video cameras are presented. The SaS mixture model is shown to improve the detection performance compared with a cluster BS method using a Gaussian mixture model and the method of Li et al. Moving object detection in video sequences represents a critical component of many modern video processing systems. The standard approach to object detection is background subtraction (BS), that attempts to build a representation of the background and detect moving objects by comparing each new frame with this representation [4]. A number of different BS techniques have been proposed in the literature and some of the popular methods include mixture of Gaussians [24], kernel density estimation [6], color and gradient cues [9], high-level region analysis [22], hidden Markov models [21], and Markov random fields [14]. Basic BS techniques detect foreground objects as the difference between two consecutive video frames, operate at pixel level, and are applicable to static backgrounds [4]. Although the generic BS method is simple to understand and implement, the disadvantages of the frame difference BS are that it does not provide a mechanism for choosing the parameters, such as the detection threshold, and it is unable to cope with multimodal distributions. One of the important techniques able to cope with multimodal background distributions and to update the detection threshold makes use of Gaussian mixture models (GMMs). The model proposed in [24] describes each pixel as a mixture of Gaussians and an online update of this model. The larger Gaussian components correspond to the background, and this is used to generate the background model. An algorithm for background modeling and BS based on Cauchy statistical distribution [13] is shown to be robust and adaptive to dynamic changes of the background scene and more cost effective than the GMM as it does not

involve any exponential operation. In [11], the foreground objects are detected in complex environments. The background appearance is characterized by principal features (spectral, spatial, and temporal) and their statistics, at each pixel. However, the learning method in [11] requires “training” since it relies on look-up tables for the features and adapts them to the changes of environment. The cluster BS-SaS technique that we propose does not need such look-up tables for the image features and is a cluster-based technique, which makes it different from [11]. According to our knowledge only one recent work [18] considers mixtures of SaS distributions for offline data analysis and does not seem suitable for real-time object detection. In this letter, we propose a novel CBS technique based on SaS distributions that we call CBS-SaS. The main contributions of the letter are threefold. Firstly, the BS process is performed at cluster level as opposed to pixel level methods that are commonly used [4], [6], [24]. The CBS-SaS method reduces significantly the clutter noise that arises owing to slight variations in the pixel intensities within regions belonging to the same object. Secondly, owing to their heavy tails, SaS distributions can help handling dynamic changes in a scene, and hence they model moving backgrounds and moving camera in a better way than the GMM. Results of modeling the background of a moving image sequence can be best obtained while operating with estimated values of the characteristic exponent parameter of the SaS distribution, rather than with fixed values corresponding to the Gaussian or Cauchy case. By estimating the parameters of the α -stable distribution, the probability density function (PDF) of clusters of pixels can be faithfully represented and a reliable model of the background can be obtained. Thirdly, we show that a mixture of SaS distributions can represent the multimodality well and guarantees reliable object detection. A wide range of tests is performed on indoor and outdoor environment, on data from a static and moving cameras.

IN “T. BOUWMANS, S. JAVED, H. ZHANG, Z. LIN, AND R. OTAZO, “ON THE APPLICATIONS OF ROBUST PCA IN IMAGE AND VIDEO PROCESSING,” PROC. OF THE IEEE, 2018.”

Robust principal component analysis (RPCA) via decomposition into low-rank plus sparse matrices offers a

powerful framework for a large variety of applications such as image processing, video processing, and 3-D computer vision. Indeed, most of the time these applications require to detect sparse outliers from the observed imagery data that can be approximated by a low-rank matrix. Moreover, most of the time experiments show that RPCA with additional spatial and/or temporal constraints often outperforms the state-of-the-art algorithms in these applications. Thus, the aim of this paper is to survey the applications of RPCA in computer vision. In the first part of this paper, we review representative image processing applications as follows: 1) low-level imaging such as image recovery and denoising, image composition, image colorization, image alignment and rectification, multifocus image, and face recognition; 2) medical imaging such as dynamic magnetic resonance imaging (MRI) for acceleration of data acquisition, background suppression, and learning of interframe motion fields; and 3) imaging for 3-D computer vision with additional depth information such as in structure from motion (SfM) and 3-D motion recovery. In the second part, we present the applications of RPCA in video processing which utilize additional spatial and temporal information compared to image processing. Specifically, we investigate video denoising and restoration, hyperspectral video, and background/foreground separation. Finally, we provide perspectives on possible future research directions and algorithmic frameworks that are suitable for these applications.

EXISTING METHOD:

The eigenvector corresponding to the minimum non-zero eigenvalue is also known as Fiedler vector and it defines two partitions of the graph based on the signs of its coefficients. We enforce the resulting sparse matrix F to act as the eigenvectors matrix of the spatial and temporal graph Laplacians. By incorporating these spectral clustering based objective functions into the low-rank decomposition, we ensure that the resulting sparse matrix encodes the spatial and temporal connectivity at the superpixel level. Incorporating spatiotemporal graph Laplacian matrices into objective function allows the proposed SSSR algorithm to detect moving objects in a

more robust manner in complex scenes, even when the appearance of the moving objects is similar to the background

PROPOSED SYSTEM:

Some RPCA enhancements have been proposed to improve the sparsity patterns of the moving objects (e.g., by [68], [78]). Zhou et al. proposed DECOLOR [78], which incorporates Markov random field constraints into the sparse matrix F . Xin et al. proposed the GFL method [68], which encodes the fused lasso regularization in the sparse component. The smoothness effect of the Markov random field and fused lasso effectively eliminates the noise and small background movements (dynamic background pixels). However, the foreground regions tend to be over-smoothed to an undesirable extent because of the strict smoothing constraints. As shown in column cin Fig.1, the over-smoothing degrades the MOD performance due to the incomplete foreground or the merging of distinct moving objects into one segment. The background region within distinct moving objects is detected as part of the moving object segment.

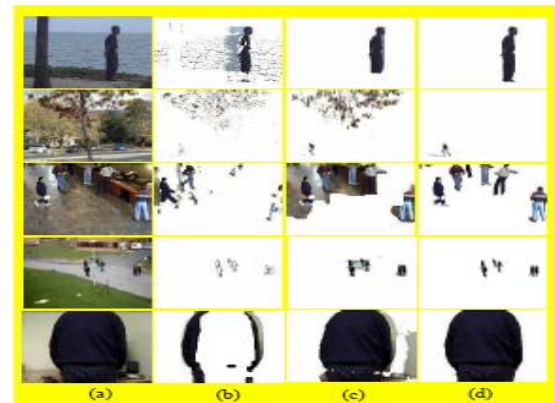


Fig. 1. Moving object detection in scenes with a complex background using RPCA methods. (a) Input images. (b) RPCA [14] results. (c) DECOLOR [78] results. (d) Results obtained using the proposed B-SSSR algorithm. From top to bottom: the first and second rows show the dynamic background 'Water Surface' and 'Fall' sequences from the I2R and *Change Detection.net* (CDnet) datasets, respectively; the third and fourth rows show the bootstrap scene 'Bootstrapping' and 'Video03' sequences from I2R and *Background Models Challenge* (BMC) datasets, respectively; and the fifth row shows 'Camouflage' from Walflower dataset, where the appearance of the moving person is very similar to the monitor in the background.

IMPLEMENTATION

ANALYSIS AND COMPARISON OF THE TWO TYPES OF MOTION DETECTION ALGORITHM

Intelligent visual surveillance system can be used many different methods for detection of moving targets, A typical method such as background subtraction method, frame difference method. These methods have advantages and disadvantages, the following will be introduced. A. Background subtraction method Background subtraction

method is a technique using the difference between the current image and background image to detect moving targets. Process flow chart is shown as Fig.

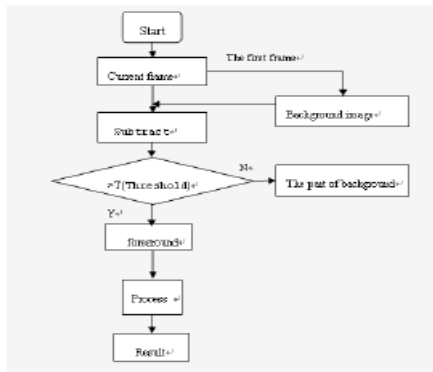


Figure 1. Flow chart of background subtraction method

The basic idea is the first frame image stored as background image. Then the current frame image k with the pre-stored background image B subtraction, And if the pixel difference is greater than the certain threshold, then it determines that the pixel to pixel on the moving target, or as the background pixel. The choice of threshold of the background subtraction to achieve the success of motion detection is very important. The threshold value is too small will produce a lot of false change points, the threshold choice is too large will reduce the scope of changes in movement. The appropriate threshold request be adapt with the impact which be had by scenes and camera on the wavelength of the color, the changes of light conditions, so the choice of the dynamic threshold should be selected [3]. The method formula is shown as (3) and (4).

$$R_k(x, y) = f_k(x, y) - B(x, y)$$

$$Dst_k(x, y) = \begin{cases} 1, & \text{background } R_k(x, y) > T \\ 0, & \text{target } R_k(x, y) \leq T \end{cases}$$

Background subtraction is used in case of the fixed cameras

to motion detection. Its advantage is easy to implement, fast, effective detection, can provide the complete feature data of the target. The shortcomings are frequent in moves of the occasions may be difficult to obtain the background image. Immovable background difference is particularly sensitive for the changes in dynamic scenes, such as indoor

lighting gradually change. The following is the video screenshot of the background subtraction method to achieve, as Fig. 2 – Fig. 5 shows



Figure 2. Background image



Figure 3. Current frame image



Figure 4. Contour map after subtraction



Figure 5. Target image

From the images we can see that a car that does not belong to the moving target appeared in the upper right corner of the target figure. This is due to the fixed background subtraction method does not process the dynamic changes in background. This is an important drawback of the method.

FRAME DIFFERENCE METHOD

Frame difference method, is also known as the adjacent frame difference method, the image sequence difference method etc. It refers to a very small time intervals Δt ($\Delta t \ll 1s$) of the two images before and after the pixel based on the time difference, and then thresholding to extract the image region of the movement, according to which changes in the region to distinguish background and moving object [4]. Frame difference of the specific flow

chart as shown in Fig. 6

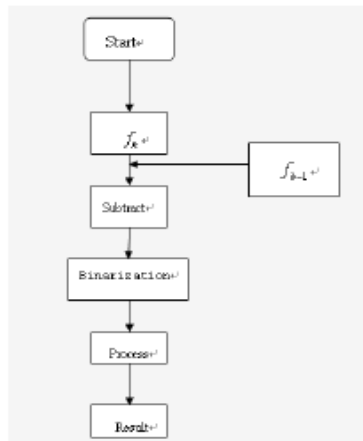


Figure 6. Flow chart of frame difference method

The specific method on calculation of difference image k Dst between the k th frame images k f with the $(k-1)$ th frame image $k-1$ f – is differential, the negative differential and fully differential, the corresponding formula is as follows:

$$\text{Differential : } Dst_k = \begin{cases} f_k - f_{k-1}, & \text{if } (f_k - f_{k-1}) > 0 \\ 0, & \text{else} \end{cases}$$

$$\text{Negative Differential : } Dst_k = \begin{cases} |f_k - f_{k-1}|, & \text{if } (f_k - f_{k-1}) < 0 \\ 0, & \text{else} \end{cases}$$

$$\text{Fully Differential : } Dst_k = |f_k - f_{k-1}|$$

The binarization for the differential image can get a collection of pixel movement. The following are the video shots of frame difference method, as Fig. 7 – Fig. 9 shows.



Figure 7. Current frame image



Figure 8. Contour map after differential



Figure 9. Target image

From the above screenshot we can see that the advantages of frame difference method is the computation of small, fast, simple, low complexity of program design. It is only sensitive to the movement of objects. In fact, only detect relative motion of the object. Because there is a very short

time interval between the two images, and the impact of the differential image by changes in light is small. So it is very suitable for dynamic changes in the scene [5]. Its drawback is that can not be completely extracted features of all relevant objects pixel point, unless the moving object itself has more complex texture features; After differential the interior of movement entities is easily empty; the non-zero area shown is generally the continuous or intermittent stripe-shaped region which is closely related with the edge of moving objects, as shown in Fig. 9. This region is more large than the region of the actual objects, its external rectangular were stretching on direction of the movement; it is very sensitive to noise and do not detect the accurate location of objects. Relative to the velocity of target, the video system sampling quickly (Δt is very small), its objectives in the location of two adjacent frames will be a very small difference. The location of the mid-point in the frame can be used as the approximate target location. If the speed of moving target detection compared with the sampling rate is very fast, this method will be improved.

MOVING TARGET DETECTION ALGORITHM BASED ON THE DYNAMIC BACKGROUND

Through the comparison of two moving target detection algorithms in the above section, in this paper it present a moving target detection algorithm based on the dynamic background. A. The dynamic update of the background In the background subtraction method, we can consider that the whole scene from two parts: the background, the foreground. Background is a static scene and which can be seen; Foreground is the moving objects which are interested in the video surveillance, such as: vehicles, pedestrians, etc [6]. However, due to the scene of the monitor changes over time, the foreground stagnation in the picture for a long time should be re-classified as part of the background; and objects which is belong to the background should be classified as part of the foreground when it starts moving. Background pixel that changes and updates over time, It is the basis of background subtraction method. In this paper, background is updated over time to re-construct the background images. The flow chart is shown in Fig.10.

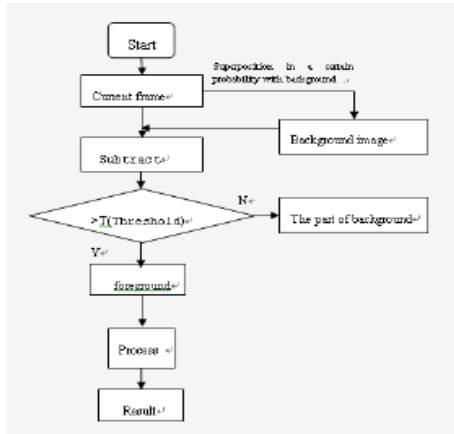


Figure10. Flow chart of moving target detection algorithm based on dynamic background

$$B_k(x, y) = B_{k-1}(x, y), f_{k-1}(x, y) \text{ superposition in a certain probability} \quad (8)$$

$$R_k(x, y) = |f_k(x, y) - B_{k-1}(x, y)| \quad (9)$$

$$Dst_k(x, y) = \begin{cases} 1, \text{background } R_k > T \\ 0, \text{target } R_k \leq T \end{cases} \quad (10)$$

The formula of the moving target detection algorithm based on the dynamic background as follows: B is the background of the kth frame image. k f is the kth frame image. The pixel in the image k B is generated from the pixel in the image k 1 f – superposition in a certain degree of probability with the pixel in the background image k 1 B – . With time, the stagnation moving targets of the video again and again as a result of superimposed to the background, in the end it can be a part into the background. And the oppositethe movement part of the background eventually separated from the background to become foreground. In this paper, the function GetBackground used to achieve background image with the current frame superposition outputting. Following introduce the used of the function GetBackground : The definition of function: GetBackground(Image*background, const Image*src_image, double alpha); Introduce of the parameters : the input image: src_image, background image: background, The weight of the input image: alpha. Function: Calculation of the input image src_image and the background image background weighted sum, and makes the image background as an average cumulative sum of the frame sequence。 The specific formula is as follows: (,) (1) (,) _ (,) background x y background x y

src image x y And α (alpha) regulates the update rate (how quickly the image background in order to forget the front of the frame). The following is the screenshots of the background image at the different time used by the new algorithm



Figure 11. Background at start time



Figure 12. After a period of time

We can see after a period of time the car at upper into the background, and the car at the right upper out of the background.

DETERMINATION OF THRESHOLD

In order to increase the adjustability of the threshold and the robustness of the background image on the brightness changes slowly. The determination of threshold as follows:

$$\theta = \left\{ mid(\theta_1, \theta_2, \theta_3, \theta_4), \theta_i = \frac{1}{N_{M_i}} \left(c \cdot \sum_{(x,y) \in M_i} R(x, y) \right) \right\}$$

And c determined by the experiment, the general admission 3-5; M_i is a region of the background, and generally selects the area at the edge; N is the area of M_i . The algorithm selected the four corners of differential gray image region to be calculated respectively, and makes the mid-value as the final check of the threshold value, and get a better result. C. Extraction of detailed images of moving targets This requires the adoption of connectedness analysis to extract the complete moving target. There are two type of connectedness: four-connected and eight-connected, as shown as Fig. 13 and Fig. 14.

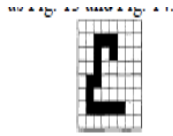


Figure 13. Four-connected

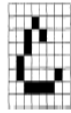


Figure 14. Eight-connected



Figure 15. Source image

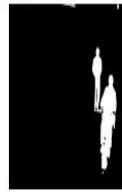


Figure 16. After connectedness
and internal be filled

$$Dst_k(x, y) = \begin{cases} f_k(x, y), & \text{if } (Dst_k(x, y)) = \emptyset \\ Dst_k(x, y), & \text{else} \end{cases}$$



Figure 17. Grayscale image



Figure 18. After connectedness
and internal be filled



Figure 19. Detailed images of moving targets

developed for moving object detection in sequences captured by moving and pan-tilt zoom cameras.

REFERENCES

- [1] R. Achanta, A. Shaji, K. Smith, A. Lucchi, P. Fua, and S. S`usstrunk, "SLIC superpixels compared to state-of-the-art superpixel methods," IEEE T-PAMI, vol. 34, no. 11, pp. 2274–2282, 2012.
- [2] H. Bhaskar, L. Mihaylova, and A. Achim, "Video foreground detection based on symmetric alpha-stable mixture models," IEEE T-CSVT, vol. 20, no. 8, pp. 1133–1138, 2010.
- [3] T. Bouwmans, S. Javed, H. Zhang, Z. Lin, and R. Otazo, "On the applications of robust pca in image and video processing," Proc. of the IEEE, 2018.
- [4] T. Bouwmans, F. El Baf, and B. Vachon, "Background modeling using mixture of gaussians for foreground detection-a survey," RPCS, vol. 1, no. 3, pp. 219–237, 2008.
- [5] T. Bouwmans, L. Maddalena, and A. Petrosino, "Scene background initialization: A taxonomy," PRL, 2017.

CONCLUSION In this study, we improved moving object detection in the presence of complex backgrounds. This is obtained by integrating the RPCA objective function and spatial and temporal graph Laplacians to constrain the structure of the sparse component. The sparse component is enforced to simultaneously contain the information in the spatial and temporal graph partitions, where each partition corresponds to the background or a moving object. The proposed objective function is solved efficiently in a batch manner using linearized ADMM and in an online manner using a novel online optimization scheme. We demonstrated that the proposed algorithms obtained excellent performance on dynamic and complex background sequences compared with 18 state-of-the-art methods using six publicly available datasets. In future, similar algorithms need to be