



TEXT EXTRACTION FROM HETEROGENEOUS IMAGES USING MATHEMATICAL MORPHOLOGY

¹ G. RAMA MOHAN BABU, ² P. SRIMAIYEE, ³ A. SRIKRISHNA

^{1,3} Dept. of CSE and IT, R.V.R& J.C College of Engg. Guntur, INDIA.

² Dept. of ECE, JNTUK, Kakinada, INDIA.

E-mail: rmbgatram@gmail.com, psrimaiyee@gmail.com, atlurisrikrishna@yahoo.com

ABSTRACT

The extraction of text in an image is a classical problem in the computer vision. Extraction involves detection, localization, tracking, extraction, enhancement and recognition of the text from the given image. However variation of text due to difference in size, style, orientation, alignment, low image contrast and complex background make the problem of automatic text extraction extremely challenging. Text extraction requires binarization which leads to loss of significant information contained in gray scale images. The images may contain noise and have complex structure which makes the extraction more difficult. This paper proposes an algorithm which is insensitive to noise, skew and text orientation. It is free from artifacts that are usually introduced by thresholding using morphological operators. Examples are presented to illustrate the performance of proposed method. The text extraction system has been attempted over a corpus of three kinds of images and promising precision has been obtained.

Keywords: *Mathematical Morphology, Morphological Operators, Edge Detection, Localization, Connected Component*

1. INTRODUCTION

Text extraction from images and video sequences finds many useful applications in document processing [1], detection of vehicle license plate, analysis of technical papers with tables, maps, charts, and electric circuits [2], identification of parts in industrial automation [3], and content-based image/video retrieval from image/video databases [4], [5]. Educational and training video and TV programs such as news contain mixed text-picture-graphics regions. Region classification is helpful in object-based compression, manipulation and accessibility. Also, text regions may carry useful information about the visual content.

However due to the variety of fonts, sizes, styles, orientations, alignment effects of uncontrolled illuminations, reflections, shadows, the distortion due to perspective projection as well as the complexity of image background, automatic localizing and extracting text is a challenging problem.

Characters in a text are of different shapes and structures. Text extraction may employ

binarization [7], [9]–[11] or directly process the original image [8], [12], [13]. In [5], a survey of existing techniques for page layout analysis is presented. Mathematical morphology is a topological and geometrical based approach for image analysis. It provides powerful tools for extracting geometrical structures and representing shapes in many applications. Morphological feature extraction techniques have been efficiently applied to character recognition and document analysis, especially if dedicated hardware is used. In this paper, we propose an algorithm for text extraction based on morphological operations.

The paper is organized as follows. In Section II, the proposed morphological text extraction technique is described. Examples and comparison with existing text extraction algorithms are presented in Section III. Conclusion is given in Section IV.

2. METHODOLOGY

The method has considered the fact that edges are reliable features of text regardless of color or intensity, layout, orientation etc. The edge

detection operation is performed using the basic operators of mathematical morphology.

Using the edges the algorithm has tried to find out text candidate connected components. These components have been labeled to identify different components of the image. Once the components have been identified, the variance is found for each connected component considering the gray levels of those components. Then the text is extracted by selecting those connected components whose variance is less than some threshold value. The complete process of text extraction is given in the form of flow chart in Figure 1.

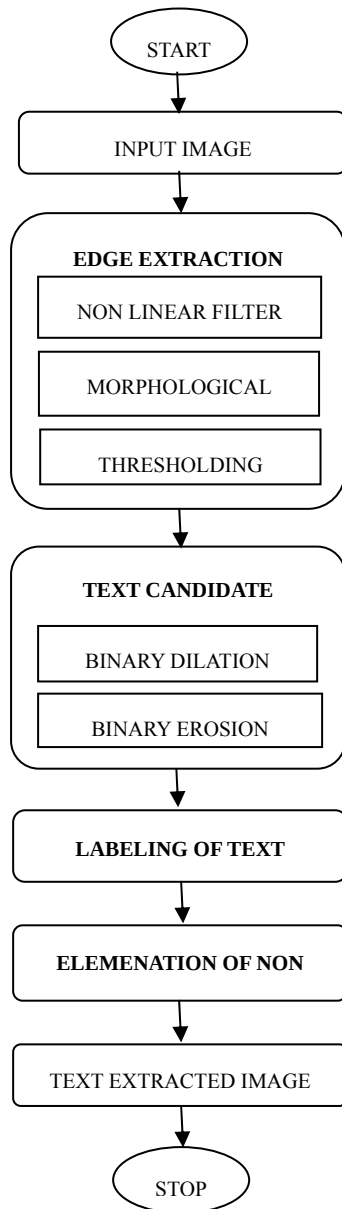


Figure 1: Flow Chart for Text Extraction.

2.1. Edge extraction

For the given input image an efficient morphological edge detection scheme is applied to find the edges of the image.

Step 1: Apply non-linear filter to the given input image to remove noise.

- In this step apply open operation to the input image.
- In this step apply close operation to the input image.
- Now find the average of the above two steps, and the resultant image is a blurred image.

Algorithm: Non-linear_Filter (y)

Input: y (original image)

Output: ybl (blurred image){

% Apply open filter on image y

$y \circ B = \delta B(\epsilon B(y))$

% δB = dilation with B on y

% ϵB = erosion with B on y

% Apply close filter on image y

$y \bullet B = \epsilon B(\delta B(y))$

% Average of the two filtered images is the blurred image

$ybl = ((y \circ B) + (y \bullet B))/2$;

}

Step 2: The blurred image obtained from the filtered image is taken as the input and we find the morphological gradient of this image.

Algorithm: Morphological gradient (ybl)

Input: ybl (blurred image)

Output: es (gradient image) {

$es = \delta B(ybl) - \epsilon B(ybl)$;

% B: 8 connected structuring element

% δB = dilation with B on ybl

% ϵB = erosion with B on ybl

}

Step 3: Threshold is used to obtain binary image from the grayscale gradient image.

The threshold value gamma is obtained using the algorithm given below.

Algorithm: Thresholding (es)

Input: es (gradient image)

Output: e (threshold image) {

$g1 = [-1 \ 0 \ 1]$

$g2 = \text{transpose}(g1)$;

$$\gamma = \frac{-\sum(es \bullet s)}{\sum s}$$

```

s = max(|g1**es| , |g2**es|)
% •denote pixel wise multiplication
% ** denotes two dimensional convolution

```

```

e = { 1 if es > γ
      0 otherwise
}

```

2.2. Text candidate region formation

From the threshold image the text candidate regions are obtained as follows.

In text candidate region formation close operation is applied to connect all the edges.

Algorithm: Region_Formation(e)

Input: e threshold image

Output: ec text candidate region images {

```
% dilation(e);
```

```
ec1 = δ s(e);
```

```
% erosion(ec1);
```

```
ec = εs(ec1);
```

```
% s is 8 connected structuring elements;
```

```
}
```

2.3. Labelling of text candidate regions

Apply labelling on the text candidate region as follows

- To the above obtained text candidate region each candidate is uniquely labelled.
- Re-labeled the text candidate regions by sequentially assigning unique values to the same component.

Algorithm: Labeling_of_Region(ec)

Input: ec text candidate region image

Output: bw labelled image {

```
bw= bwlabel(ec);
```

```
% calculates all the connected components with sequential label;
```

```
}
```

2.4. Elimination of non text region

From the labelled image which contains text and non text regions eliminate the non text regions using variance operation as follows

Algorithm: Region_Elimination(bw)

Input: bw labelled image

Output: ext image containing text {

$$var = \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2$$

```
% Finds variance on the original Gray scale image on the labelled components
```

```
% Select the regions whose variance values are less than threshold.
```

```
% Threshold is calculated by taking average of all gray level values
```

```
ext = union( var(images) < threshold );
```

```
}
```

3. Experimental Results and Discussion

3.1. Experimentation Setup

This algorithm has been tested over a corpus of 60 text images of three type including caption text, scene text and document images in which text has different font size, color, orientation, alignments . These images are analyzed to demonstrate the performance of the proposed algorithm. Performance is verified with the oriented text in horizontal & vertical direction with different languages (English & Hindi). Various metrics have been evaluated from the tested results.

3.2. Performance Analysis

Metrics used to evaluate the performance of the system are Precision, Recall and F-Score. Precision and Recall rates have been computed based on the number of Correctly Detected Characters (CDC) in an image, in order to evaluate the efficiency and robustness of the algorithm. The metrics are as follows:

Definition 1: False Positives (FP) / False alarms are those regions in the image which are actually not characters of a text, but have been detected by the algorithm as text.

Definition 2: False Negatives (FN)/ Misses are those regions in the image which are actually text characters, but have not been detected by the algorithm.

Definition 3: Precision rate (P) is defined as the ratio of correctly detected characters to the sum of correctly detected characters plus false positives as represented in equation below.

$$P = \frac{CDC(true\ positives)}{CDC + FP} \times 100\%$$

Definition 4: Recall rate (R) is defined as the ratio of the correctly detected characters to sum of correctly detected characters plus false negatives as represented in equation below.

Definition 5: F-score is the harmonic mean of recall and precision rates as represented in equation below.

3.3. Results and discussion

The algorithm is tested on the oriented text in horizontal & vertical direction. The output images of the proposed algorithm in Figure 2, 3 and 4 consist of detected text for caption text, document images and scene text respectively.

The results obtained by the proposed algorithm on three types of sample images are presented in Table I where the number of True positives, number of False alarms/False positives, number of False negatives/Misses and the corresponding values for Precision, Recall and F-score are listed. Averages of all measures for three types of images have been listed in Table II.



Figure 2: Text extraction for caption text images



Figure 3: Text extraction for document images



Figure 4: Text extraction for scene images

TABLE I

PERFORMANCE MEASURES FOR THE PROPOSED METHOD (FOR A DATA SET WITH SAMPLE IMAGES)

Image Type	Image name	TP/ Detections	FP/ False Alarms	FN/ Misses	P	R	F-SCORE
Document Text Image	SS5	13	55	4	79.16	76.47	77.79
	DOCM8	13	3	2	81.25	86.66	83.86
	DOCM11	47	5	7	90.38	90.38	90.38
Scene Text Image	S2	12	5	0	70.58	100	82.75
	CC2	64	22	13	74.41	83.11	78.51
	S21	18	30	0	37.5	100	41.85
Caption Text Image	CAPTION3	12	1	2	92.30	85.71	88.88
	CC1	5	1	1	83.33	83.33	83.33
	CAPTION2	23	14	2	62.16	92.00	74.19

TABLE II

AVERAGE PERFORMANCE MEASURE FOR THE PROPOSED METHOD

Image sets Measures	Caption Text Images	Scene Text Images	Document Images
Number of True Positives (TP)	13.33	32	24
Number of False Positives (FP)	7.33	19	21
Number of False Negatives (FN)	9.22	5	8
Recall	87.01	94.37	84.53
Precision	79.26	60.86	83.59
F-measure	82.33	67.70	84.01

After experimenting it has been found that the Precision rate & intern average F-score of scene text images dropped down when compared to caption text images & document images, because of their embedded nature of text inside the image. The result shown in the table is compared by means of bar chart which shows that the caption and document images have good precision than that of scene images. The precision rate vs. recall rate is shown in the Figure 5.

3.4. COMPARISON WITH OTHER TEXT EXTRACTION TECHNIQUES

To give an average estimate of the performance of the text extraction the results have been compared against two existing algorithms [14] and [15]. Both the methods have used the aspect ratio to identify the text and non text regions within an image. The first method has used the complicated procedure of finding inner, outer and inner-outer corners.

The second procedure has identified edge at different orientation i.e. 0, 45, 90, and 135 degrees and grouping these strokes at different heights, text is extracted. This increases the complexity of algorithm to identify the edges at different orientations. The new Connected Component Variance (CCV) approach has solved the above problems.

In order to group the isolated characters to a meaning full word, a dilation operation with varying structuring element is used. The algorithm identifies number of components and calculates the variance of each component if the variance of each component is high then it is believed to a kind of symbol rather than a text.

This algorithm is in sensitive to skew and text orientation, the output of the text extraction algorithm is fed to an OCR system to recognize the contained information. The main objective of the text extraction algorithm is to reduce the number of false text candidate that may be fed to the OCR.

Further investigation on the threshold value to select the correct text candidate is being performed. Different images with different lighting and contrast is used for text extraction. These images are tested using the two algorithms Threshold with Average Gray values (TAG) and the other Threshold Using Morphological operators (TMO). To evaluate the performance of TAG and TMO 25 text images with different font size, perspective, alignments any number of characters in a text string under different lighting

conditions is considered. It has been observed that in many case the images gives good text extraction when TMO is used, when the image has uniform background. If various foregrounds and backgrounds are presented in the image, it has been observed that TMO is in sensitive to noise and introduce minimum noise on the removal of non text information. The TAG algorithm introduces noises while removing non text characters but has an advantage of extracting those text characters whose gray levels are close to the back ground shown in the Figure 6.

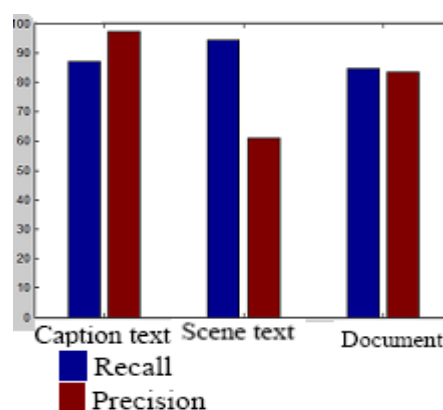


Figure 5: Precision vs Recall rate graph.



Figure 6: Text extraction using TMO and TAG algorithms.

4. CONCLUSION

This paper proposes a new text extraction algorithm from a text/graphics mixed document images. This algorithm is in sensitive to skew and text orientation the output of the text extraction algorithm is fed to an OCR system to recognize the contained information. The results obtained on varied set of images are compared with respect to precision and recall rates. Promising results have been obtained on a number of images in which almost all text lines can be retrieved from the graphics and figure regions. The images have been tested using various threshold techniques. It has



been observed that the threshold depends on a various parameters like the illumination condition, reflections and the scan point spread function. This approach has used morphological clearing of images which would help to reduce the number of false positive obtained. This cleaning of the image could result in a higher precision rate.

Further investigation on the threshold value to select the correct text candidate is being performed. The images have been tested using various threshold techniques. It has been found that TAG gives efficient extraction comparative to TMO when the images are taken in poor illumination. However the TAG gives noisy distorted binary images comparative to TMO. This approach has used morphological clearing of images which would help to reduce the number of false positive obtained. In corporately the OCR algorithm with proposed morphological text extraction method yields a useful system for text analysis in images.

REFERENCES

- [1] K. Jain and B. Yu, "Document representation and its application to page decomposition," *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 20, pp. 294–308, Mar. 1998.
- [2] K. Jain, and B. Yu, "Automatic Text Location in Images and Video Frames", *Pattern Recognition*, 31 (12) (1998) 2055-2076.
- [3] K. Jain, and Y. Zhong, "Page Segmentation using Texture Analysis", *Pattern Recognition*, 29 (5) (1996) 743-770.
- [4] Chitrakala Gopalan and D. Manjula, "Contourlet Based Approach for Text Identification and Extraction from Heterogeneous Textual Images", *International Journal of Computer Science and Engineering* 2(4) (2008) pp.202-211.
- [5] D. Crandall, S. Antani, and R. Kasturi, "Robust Detection of Stylized Text Events in Digital Video", *Proceedings of International Conference on Document Analysis and Recognition*, 2001, pp. 865-869.
- [6] J. Serra, *Image Analysis and Mathematical Morphology*. New York:Academic, 1982.
- [7] K. C. Fan, L. S. Wang, and Y. K. Wang, "Page segmentation and identification for intelligent signal processing," *Signal Process.*, vol. 45, pp.329–346, 1995.
- [8] R. Lienhart, "Indexing and retrieval of digital video sequences based on automatic text recognition," in *Proc. ACM Int. Conf.*, Boston, MA, Nov. 1996.
- [9] K. Jain and B. Yu, "Document representation and its application to page decomposition," *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 20, pp. 294–308, Mar. 1998.
- [10] J. Ohya, A. Shio, and S. Akamatsu, "Recognizing characters in scene images," *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 16, pp. 215–220, Feb. 1994.
- [11] H. M. Suen and J. F.Wang, "Text string extraction from images of colorprinted documents," *Proc. Inst. Elect. Eng. Vis., Image, Signal Process.*, vol. 143, no. 4, pp. 210–216, 1996.
- [12] L. Wang and T. Pavlidis, "Direct gray-scale extraction of features for character recognition," *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 15, pp. 1053–1067, Oct. 1993.
- [13] Y. Zhong, K. Karu, and A. K. Jain, "Locating text in complex color images," *Pattern Recognit.*, vol. 28, no. 10, pp. 1523–1535, 1995.
- [14] Jagath Samarabandu, Member, IEEE, and Xiaoqing Liu 2007 "An Edge-based Text Region Extraction Algorithm for Indoor Mobile Robot Navigation" ,*International Journal of Signal Processing* 3(4)2007.
- [15] Yassin M. Y. Hasan and Lina J. Karam 2000 "Morphological Text Extraction from Images" *IEEE Transaction on Image Processing* vol. 9,no. 11,Nov. 2000.

AUTHOR PROFILES:

G. Rama Mohan Babu, received his B.Tech degree in Electronics & Communications Engineering from Sri Venkateswara University, India. He did his M.Tech in Computer Science & Engineering from Jawaharlal Nehru Technological University (JNTU), India. He is currently working as Associate Professor, in the Department of Information Technology at RVR & JC College of Engineering, Guntur, India. He has 12 years of teaching experience. He is pursuing his Ph.D from Acharya Nagarjuna University (ANU), India, in Computer Science & Engineering under the guidance of Dr. B. Raveendra Babu. His research areas of interest include Image Processing, and Pattern Recognition. He is life member in professional bodies like ISTE and CSI.



P. Srimaiyee is pursuing her Bachelor degree in Electronics & Communications Engineering, from Jawaharlal Nehru Technological University Kakinada (JNTUK) Kakinada, India. From the inception of the course she has been involved in various innovative projects. She has won many awards in technical quiz and has actively presented papers and participated in the student technical symposiums and seminars at national level.



A. Srikrishna received the AMIE degree in Electronics & Communication Engineering from Institution of Engineers; Kolkatta in 1990, M.S degree in Software Systems from Birla Institute of Technology and Science, Pilani in 1994, M.Tech degree in Computer Science from Jawaharlal Nehru Technological University (JNTU) in 2003. She has worked as lecturer from (1991-95). She is working in RVR & JC College of Engineering, Guntur from 13 years as Assistant Professor and Associate Professor. She is pursuing her Ph.D from Jawaharlal Nehru Technological University (JNTU) in Computer Science under the guidance of Dr.V.Vijaya Kumar. Her research interest includes Image Processing and Pattern Recognition. She is Associate member of IE (I) and member of CSI.