

A Method for Extraction of Text from Image file and Number Plate of Vehicle

Abhijeet Singh Thakur^{1*} Neelesh Mehra²

¹ M. Tech. Student, Samrat Ashok Technological Institute, Vidisha (MP)

² Assistant Professor, Samrat Ashok Technological Institute, Vidisha (MP)

Abstract – In the computer vision research area image text extraction and extract number from vehicle number plate is a well-known problem. Planned picture details, organizing of pictures, content-based data indexing and recovery depend on the textual data demonstration in those images. Its challenging and difficult job to take away content from images due to the variations in the content, for example, content scripts, style, text style, size, shading, arrangement and orientation; and because of external factors, for example, low picture differentiate (textual) and complex background.

Investigational outcomes on different dataset illustration that proposed scheme leads to a high performance.

Keywords—Extraction of text, text localization, text recognition, Image Pre-processing.

-----X-----

I. INTRODUCTION

Nowadays, information libraries that primarily limited untainted content are ending up gradually boosted by collaborating media parts, for instance, images, recordings and sound clasps. These parts want a planned aims to well index and focus interactive media parts. At the point when content display in images it could be known, separated, and observed subsequently, they would be an important wellspring of high level semantics.

Text Extraction are often explain as text extract from pictures, extracting the related text information from input images. Increasingly progress of digital media has resulted in digitization of all classes of materials. Lot of resources is offered in electronic medium. The electronic medium converted to pictures. This kind of pictures shows many difficult investigation issues in text extraction and recognition. Analysis of Document, detection of vehicle license plate, Analysis of article with tables, maps, graphs, charts then on., watchword primarily based image look, recognizable proof of elements in mechanical mechanization, content primarily based retrieval, name plates, identification of object, road signs, content primarily based video classification, video content extraction, segmentation of page, report retrieving, address block space then on are applications of text extraction from pictures.

Images can be generally categorized into 3 type:

- Document images
- Caption text images
- Scene text images.

Figures 1-4 demonstrate a number of cases of content in photos. A document image (Figure 1a, 1b) additional usually than not contains content and few graphics segments. Scan the document photos are obtained by examining printed document, journal, manually written historical archive, corrupted document photos, and book cover so on. Free content from image it would show up during a for all intents and purposes boundless variety of text designs, style, arrangement, estimate, shapes, colors, and so on. Content extraction from document with content on complex shading background is difficult due to varied nature of the background and stir up of color(s) of fore- ground content with shades of background.

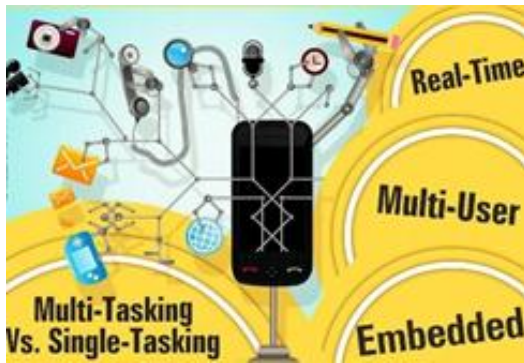


Figure 1a: Document Text Image Figure



Figure 1b: Colored Text Image

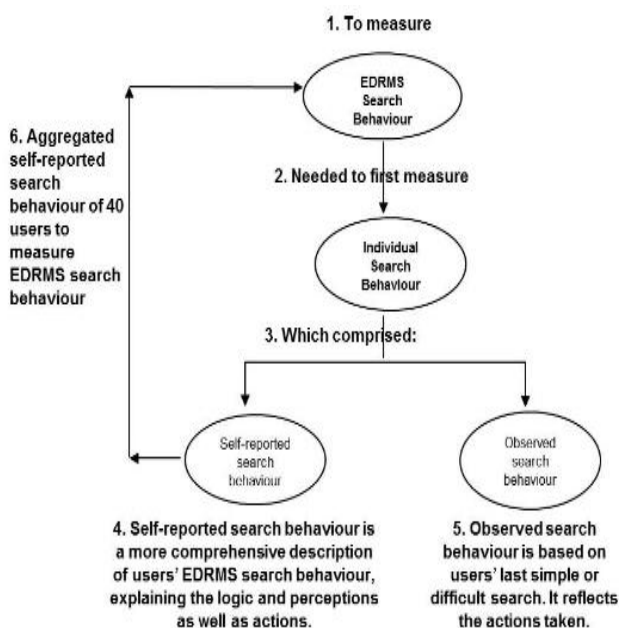


Figure 2: Caption Text Image



Figure 3: Scene Text Image



Figure 4: vehicle Text Image

II. PROCESS IN TEXT RECOGNITION

Preprocessing is that the initial phase in the handling of scanned image. The scanned image is checked for noise, skew, incline so on. Image obtaining skew with either left or right orientation or with noise, for instance, Gaussian these are potential outcomes comes at the time of extraction. Here the image is initial modification over into grayscale and afterward convert into binary image. Consequently we tend to get image that is suitable for in addition process. When pre-preparing, the noise is expelled from image presently passed it to segmentation stage, during this stage image is isolated into singular characters. The binary image is tested for inter line areas. Within the event that inter line areas are recognized then the image is portioned into sets of sections over the interline gap. Under sections lines are scanned for even area crossing purpose as for the background. Histogram of the image is employed to recognize the width of the horizontal lines. For vertical area convergence lines are scanned vertically. Here histograms are used to spot the width of the words. When this word are separated into characters utilizing character width calculation technique. Segmentation amount of OCR takes when by include extraction wherever the individual image glyph is considered and extricated for highlights. Initial a character glyph is characterized by the accompanying qualities like stature of the character, width of the character. Classification is finished utilizing the highlights extricated within the past advance, which compares to every character glyph. These highlights are investigated utilizing the arrangement of principles and marked as having an area with varied categories. This classification is summed up with the top goal that it works for single text style composes. The stature of the character and also the width of the character, totally different separation measurements are picked because the risk for classification once strife happens. Thus additionally the classification rules are composed for various characters. This strategy may be a nonexclusive one since it removes the state of the characters and want not be prepared. At the purpose once another glyph is given to the current classifier block it extricates the options and compares regarding the options in line with the principles and afterward perceives the character and marks it.

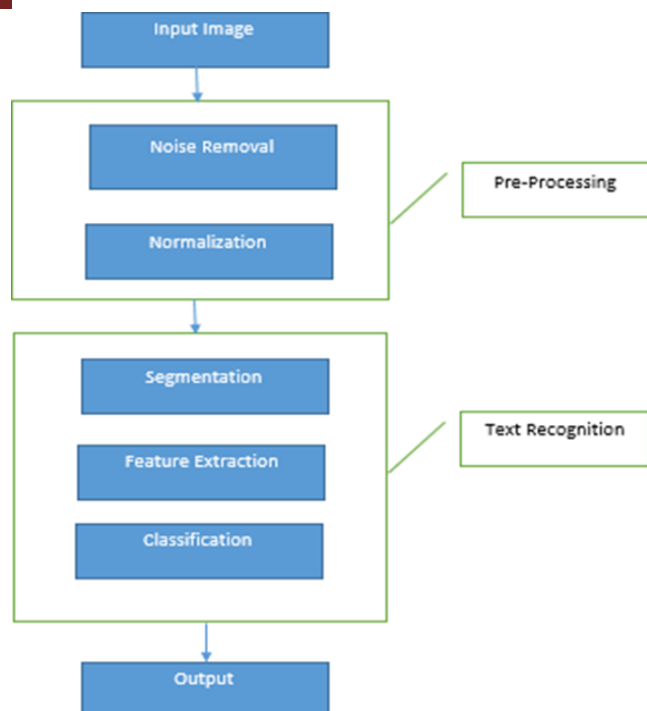


Figure 5: Process in text recognition

ALGORITHMS

1. Start
2. Scan the textual image.
3. The next step is Color image convert into gray image and then binary image.
4. After done preprocessing procedure like noise removal, skew correction etc.
5. Load the DATABASE.
6. Do segmentation by separating lines from textual image.

III. PROPOSED METHODOLOGY

We planned a work for acknowledgment and retrieval tasks within the large dictionary setting. we tend to distinguish potential character areas and find words contained within the image. we tend to reduce the big lexicon to a little image specific dictionary. The vocabulary decrease method exchanges between re-registering priors and refining the dictionary. we tend to assessed planned work results on public datasets and show predominant execution on giant and medium lexicons for acknowledgment and recovery tasks.

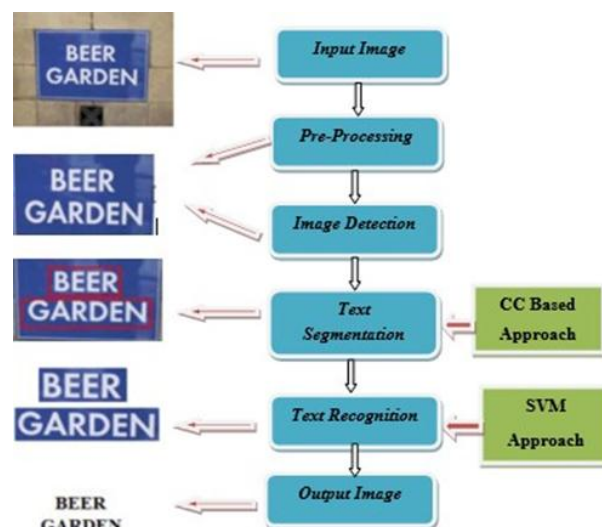


Figure: 6 Flowchart of a Proposed Work

IV. IMPLEMENTATION

This work is meant to handle numerous types of pictures effectively. So, the image process programs for numerous methodologies during this thesis are implemented using Matlab toolbox (Image processing).

Table: 1 Dataset details for Implementation

S. No	Image Type	No. Of Images
1	Image Text	30
2	Document Text	30
3	Vehicle Number Plate	40

Proposed methodology has tested about one hundred samples of Image Text with 30 pictures, Document Text with 30 pictures, Vehicle number Plate with 40 pictures all of them are over one thousand characters that are manually extracted the from the Matlab2014a. We have not actual variety of characters due to documented pictures. Some image text carry 10-15 characters, in documented text they are in paragraph, in vehicle variety plate image some plate variety are 7 in total and a few are 10 in totals we offer section wise result. Below table shows accuracy of total recognized pictures with 73%.



Niels Bohr defined a profound truth as a truth whose opposite is also a (profound) truth. In that spirit, I'd like to define a deep question as a question that doesn't make clear sense until after you've answered it. The election problem is a deep question. In posing it, we take it for granted that we can make a distinction between what is "possible" and what is "real." On the face of it, that sounds pretty unscientific: the goal of science is to understand the real world, and only the real world is possible! But in practice a clear and useful distinction does emerge.



Figure 7: Data set for Text image, document image, vehicle number plate.

We have taken random data set for text image and documented dataset some are books cover page, some are company's logo and banners for brand. For documented dataset we have taken scanned notes, long notice etc. In case of vehicle number plate we have taken Ireland vehicle number plate data as well some Indian number plate data.

Table-2 Implementation of Text image, document image, vehicle number plate

S. No	Input image	Detect image
1		
2		
3		
4		
5		

Table 3 Accuracy of Total Recognized Images

Total Images	Total Unrecognized Images	Total Recognized Images	Accuracy (%)
100	27	100-27=73	73%

The investigation was administered on varied pictures. The performance of algorithmic rule has been estimated supported its preciseness rate and recall rate. True positives are the non-text regions within the image and are detected by the algorithmic rule as text regions. True negatives are the text regions in the image and have not been detected by the algorithm. Both rate of recall and rate of precision are suitable as measures to find the accuracy and sensitivity of planned work in removing the non-text regions and placing the correct text regions. Evaluate the performance of proposed method is using sensitivity and specificity.

Confusion Matrix

A confusion matrix may be a table that's usually used to describe the performance of a classification model on a collection of check information that the true values are best-known.

Table 4: Confusion Matrix

	P (Predict)	N (Predict)
P(Actual)	True Positive (T.P.)	False Negative (F.N.)
N(Actual)	False Positive	True

TN / True Negative: case was negative and predicted negative

TP / True Positive: case was positive and predicted positive

FN / False Negative: case was positive but predicted negative

FP / False Positive: case was negative but predicted positive

Sensitivity (precision rate or TPR) = TP / TP

+ FN

Sensitivity = $73/73+1=98\%$

Specificity (recall rate or FPR) = TN / TN + FP

Specificity= $7/7+1=87\%$

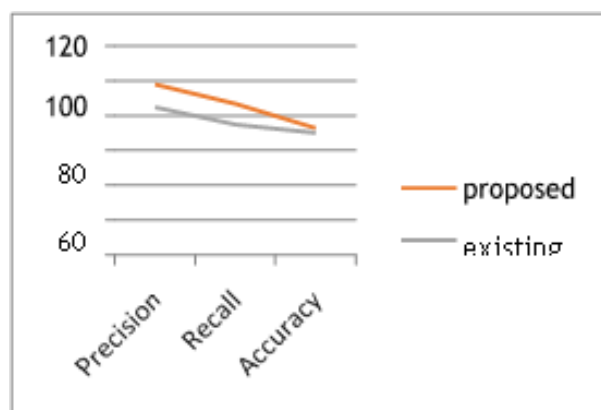


Figure 8: Comparison performance between proposed and existing work

Precision is that the ratio of correctly expected positive observations to the whole expected positive observations. High preciseness relates to the low false positive rate. Recall is that the ratio of properly expected positive observations to the all observations in actual category. Accuracy is that the most spontaneous performance measure and it is basically a ratio of correctly expected observation to the whole observations. (Rajan & Raj, 2017) planned new fractional Poisson model for improving fine specifics in natural scene images by considering the noises formed by Laplacian operation, scholars simply work on scene image but proposed work on text image, documented image, vehicle number plate. As compare to author's work proposed work perform better.

V. CONCLUSIONS

Content objects happening in image and video reports will offer a lot of valuable information to content primarily based information recovery and categorization applications, since they contain a lot of semantic information known with the records' substance. Be that because it could, separating content from pictures and videos is an exceptionally troublesome part due to the shifting textual style, estimate, shading, introduction, and distortion of content articles. In spite of the actual fact that countless extraction approaches are accounted for within the writing, no particular composed content model and character options are introduced to catch the novel properties and structure of characters and content articles.

This paper targets the examination and advancement of text identification and text acknowledgment calculations for content in planned digital photos, vehicle range plate and natural scene photos.

In this paper, we tend to project a unique unsupervised strategy to spot and localize the

content things happening in image and documents. For content localization and recognition, by expecting that a content protest contains in excess of a thousand characters with a hundred photos and preciseness with 73, we tend to build a text model visible of pictorial structure.

The exploratory outcomes demonstrate that the projected techniques will accomplish most well-liked binarisation results over traditional binarisation ways. Each subjective and quantitative assessment demonstrates the viability of the projected techniques.

VI. FUTURE WORK

The future work for the most part focuses on build up a calculation for proper and fast content extraction from an image. This application is conceivable to vary over a discourse to content in addition. The calculation has thus far been tried on content prevailing documents simply that are checked from daily papers, magazines, and story books of children, postal envelopes. In addition we are going to attempt the projected approach apply on synthesized pictures.

VI. REFERENCES

1. B. Epshtein, E. Ofek, and Y. Wexler (2010). Detecting text in natural scenes with stroke width transform. In CVPR.
2. L. Kang, Y. Li, and D. Doermann (2014). Orientation robust text line detection in natural images. In CVPR, 2014.
3. L. Neumann and J. Matas. Text localization in real-world images using efficiently pruned exhaustive search. In ICDAR, 2011.
4. L. Neumann and J. Matas. Real-time scene text localization and recognition. In CVPR, 2012
5. L. Neumann and J. Matas (2013). Scene text localization and recognition with oriented stroke detection. In ICCV, 2013.
6. C. Shi, C. Wang, B. Xiao, Y. Zhang, and S. Gao (2013). Scene text detection using graph model built upon maximally stable external regions. Pattern Recognition Letters, 2013.
7. L. Gomez and D. Karatzas (2013). Multi-script text extraction from natural scenes. In ICDAR.

8. L. Gomez and D. Karatzas (2014). A Fast Hierarchical Method for Multi-script and Arbitrary Oriented Scene Text Extraction.
9. M. Jaderberg, K. Simonyan, A. Vedaldi, and A. Zisserman (2014). Reading text in the wild with convolutional neural networks. *International Journal of Computer Vision*.
10. M. Jaderberg, A. Vedaldi, and A. Zisserman (2014). Deep features for text spotting. In *ECCV*.
11. Ming Zhao, Shutao Li, and James Kwok (2010). Text detection in images using sparse representation with discriminative dictionaries. *Image Vision Comput.*, 28(12): 1590–1599, December 2010. ISSN 0262-8856.
12. De Campos, T., Babu, B., and Varma, M. (2009). Character recognition in natural images. In *International Conference on Computer Vision Theory and Applications*.
13. Yi, Chucai, Yang, Xiaodong, and Tian, Yingli (2013). Feature representations for scene text character recognition: A comparative study. In *International Conference on Document Analysis and Recognition*.
14. Tian, Shangxuan, Lu, Shijian, Su, Bolan, and Tan, Chew Lim (2013). Scene text recognition using co-occurrence of histogram of oriented gradients. In *International Conference on Document Analysis and Recognition*.
15. M. R. Lyu, Jiqiang Song, and Min Cai (2005). A comprehensive method for multilingual video text detection, localization, and extraction. *IEEE Trans. Cir. and Sys. for Video Technol.*, 15(2):243–255, February 2005. ISSN 1051- 8215.
16. KwangIn Kim, Keechul Jung, and Jin Hyung Kim (2003). Texture-based approach for text detection in images using support vector machines and continuously adaptive mean shift algorithm. *IEEE Trans. Pattern Anal. Mach. Intell.*, 25(12): pp. 1631–1639, December 2003.
17. Huizhong Chen, Sam S. Tsai, Georg Schroth, David M. Chen, Radek Grzeszczuk, and Bernd Girod (2011). Robust text detection in natural images with edge-enhanced maximally stable extremal regions. In 2011 IEEE.
18. Lluís Gomez and Dimosthenis Karatzas (2014). Mser-based real-time text detection and tracking. In 22nd International Conference on Pattern Recognition.
19. Tao Wang, David J. Wu, Adam Coates, and Andrew Y. (2014). Ng. End-to-end text recognition with convolutional neural networks.
20. ZohraSaïdane and Christophe Garcia (2007). Automatic Scene Text Recognition using a Convolutional Neural Network. In *International Workshop on Camera-Based Document Analysis and Recognition (CBDAR 2007)*, pages 100–106, September 2007. URL <http://liris.cnrs.fr/publis/?id=6094>.
21. C. Strouthopoulos and N. Papamarkos (1998). Text identification for document image analysis using a neural network. *Image and Vision Computing*, 16 (12–13):879 – 896, ISSN 0262-8856.
22. Kai Wang and Serge Belongie (2010). Word spotting in the wild. In *Proceedings of the 11th European Conference on Computer Vision: Part I, ECCV'10*, pages 591–604, Berlin, Heidelberg, 2010.
23. ShangxuanTian, Shijian Lu, Bolan Su, and Chew Lim Tan (2014). Scene text recognition using co-occurrence of histogram of oriented gradients.
24. T. E. de Campos, B. R. Babu, and M. Varma (2009). Character recognition in natural images. In *Proceedings of the International Conference on Computer Vision Theory and Applications*, Lisbon, Portugal, February 2009.
25. M. Jaderberg, A. Vedaldi, and A. Zisserman. Deep features for text spotting. In *European Conference on Computer Vision*, 2014.
26. Veena Rajan, Shani Raj (2017). “Text Detection and Character Extraction in Natural Scene Images Using Fractional Poisson Model” *Proceedings of the IEEE 2017 International Conference on Computing Methodologies and Communication (ICCMC)* 978-1-5090-4890-8/17/\$31.00 ©2017 IEEE.

Corresponding Author

Abhijeet Singh Thakur*

M. Tech. Student, Samrat Ashok Technological
Institute, Vidisha (MP)

thakurbittu93@gmail.com