

An Approach Script Identification from Multilingual Documents using KNN and CNN

Thenarasi V¹ and Sharath Kumar Y H²

¹Government First Grade College, Siddhartha Nagar, Mysore
Email: Thenu.sunitha@gmail.com

²Department of ISE, Maharaja Institute of Technology Mysore
Email: sharathyhk@gmail.com

Abstract— Abstract--In the context of our country, script recognition is a complex task because of more number of prevailing scripts and that many of the official documents are multiscripts in nature, English is being commonly used along with regional script. Hence, it is pre-requisite to identify the script in the document and then feed the document to the respective OCR for further processing. In this work, we explore the significance of Gabor and SIFT operator for script identification at the word level respectively. Identification of script type is done by using K-NN and CNN classifiers.

I. INTRODUCTION

Now a days, handwritten text recognition is posing a lot of research openings as it improves the interface between a man and a machine in various applications, such as automation process in post offices, number plate recognition, processing of historical documents, bank cheque processing, etc,. Issues related to the printed documents are attempted satisfactorily by the research community, where as an analysis of handwritten documents is still a very challenging issue because of handwriting style varies from an individual to an individual, with size, shape and orientation of the characters. Although several pieces of research exist on monolingual handwritten text processing, negligible amount of work has been reported towards multilingual handwritten text processing. Our nation, India, is a multilingual country where many official documents contain at least three languages; one an official language of the local state, national language Hindi and the third language English, which is common for the purpose of official communication. Thus, handwritten multilingual documents pose much challenging issues in document analysis and recognition process when compared to a monolingual document, a document with a single script. An important challenge in handwritten documents is tackling the inherent skew introduced due to the writer's handwriting. Complexities such as variable spacing between words and lines, variable line skew, variable line width and height, overlapping words and lines etc., arise in handwritten documents. Thus estimating multiple skews in a handwritten document is essential in document processing.

II. RELATED WORKS

Research in script identification is not a new one. Past many years, methods have been proposed for script identification from scanned document images. More research is done in identification of script from printed documents compared to research in script identification from handwritten documents. Handwritten documents are

common in developing countries like India. Moreover, many scripts are used in India compared to any other country in the world [7]. Given a document image, script identification task can be carried out at line level, word level and block level. Block level script identification identifies the script of the given document in a mixture of various script documents. If document appears in single script, a block of the document is sufficient to identify the script type [8]. Hence, input document to the block level identification is a mono script document. In case of multi-script document, line level or word level identification is preferred. Many of the reported studies in the literature accomplish script recognition either at the line level or at the word level. In line based script identification, a document image can contain more than one script but it requires the same script on a single line. Word level script identification allows the document to contain more than one script and the script of every word is identified. However, these components are applicable only after line and word segmentation of the underlying document image. However, such a fine segmentation is not required at block level and consequently simplifies the script classification task. Generally, script identification task comprises of three stages viz, pre- processing, feature extraction and recognition. Pre-processing eliminates unwanted objects in the document image (noise, background, etc.) and extracts the words or lines or blocks from the image for feature -extraction. The feature extraction method extracts the relevant data (usually a reduced representation of the input object), called features vector, which is then fed to a classifier for script recognition. A brief overview of different methods available in the Literature for script identification is presented below. Emphasis is given on the recognition methods for Indian script as domain of our study confines to Indian scripts. Using fractal based features, S. Ben Moussa et al., in [9] have reported script identification for Arabic and Roman scripts at line level of the text document. Profile based script identification technique is presented by M.C. Padma et al., in [10]. Distinct features at top and bottom of the text line are considered for feature extraction. Features are extracted by calculating density ratio, pixel distribution, max pixel density at the bottom row of text line and density of connected components at the top of the text line. Classification is performed by using KNN. Mallikarjun Hangarge et al., in [11] have proposed a visual texture based methodology for handwritten script recognition at line level and block level. Horizontal projection profile is used to extract non-touching text line from document image. Morphological filters are utilized to extract 13 spatial spread features. KNN classifier has been used to classify the script. Rajput et al., in [12] presented line level script recognition using Gabor filter combined with DCT and wavelets. Nine Indian scripts were used for performing experiment, CNN classifier reported better recognition accuracy compared to the performance by KNN classifier. Prasanthkumar et al., in [13] have proposed an automatic separation of text line followed by word segmentation and computed morphological and Gabor features. Classification was carried out by using MLP, CNN and KNN classifiers with the observation that MLP classifier yielded better results over CNN and KNN. Long Short Term Memory (LSTM) architecture based script identification is reported by Adnan Ul-Hasan in [14]. DCT and distance transform based script identification from handwritten document is proposed by S.K. Md. Obaidullah et al., in [15]. Bangla, Roman, Devanagari and Oriya scripts are considered for performing experiments. Feature vector based on texture (BRT, BDCT, BFFT and BDT) are used to classify bi-script and tri-scripts. Gabor filter and morphological re-construction is proposed by Nibaran Das et al., in [16] to identify the Indic scripts by using MLP classifier. B. Shi et al., in [17] have incorporated deep learning based script identification from natural images. They have presented a two stage features extraction methodology i.e., feature extraction and discriminative clustering at mid-level representation, global fine tuning at second stage by modeling feature extraction and classification into one neural network by transferring learned parameters at first stage. Back propagation is used to train the network. Script dependent and script independent feature based classification of script at block level, line level and word level is proposed by S.K. Obaidullah in [18]. Classification of script is carried out by using MLP and random forest classifiers. Md. Obaidullah et al., in [19] have proposed a method for handwritten script identification based on structural, directional and texture based features. The features were extracted from document image in binary and input to MLP for text identification. Judith Hochberg et al., in [20] have reported recognition of six Indic and non-Indic scripts, viz, Arabic, Chinese, Cyrillic, Devanagari, Japanese, and Latin, using features like horizontal and vertical centroids, sphericity, aspect ratio, white holes, and so forth. The experiments have been carried out at document level. V Singhal et al., in [21] proposed a technique to identify four Indic scripts, viz, Devanagari, Bangla, Telugu, and Latin. They have used rotation invariant texture features using multi-channel Gabor filtering and gray level co-occurrence matrix. It is observed that many of the approaches are sensitive to scaling and rotation. Very recently, Scale Invariant Feature Transform (SIFT) technique has been proved to be a promising image features, for object recognition, that provides a set of features of an object that are not affected by many of the complications experienced in other methods including object scaling and rotation proposed by David G. Lowe in [22]. Recognition of Bangla and English scripts has been presented by Lijun Zhou et al., in [23] using connected component profile based features. Experiments are performed on the document text at

line, word and character level. Multi-script identification at word level based upon features extracted using Discrete Cosine Transform (DCT) and Wavelets of Daubechies family have been proposed by G .G Rajput et al., in [24]. The recognition has been carried out by using neural network for handwritten documents of nine Indian scripts including English script. K. Roy et al., in [25] have reported recognition of six popular Indic-scripts, viz, Bangla, Devanagari, Malayalam, Urdu, Oriya and Roman by using component based features, fractal dimension based features, circularity based features and so forth. This is the first kind of work involving six Indic scripts altogether. Mallikarjun Hangarge et al., in [26] have proposed a word level scheme based on directional DCT based feature to recognize six Indic-scripts, viz, Roman, Devanagari, Kannada, Telugu, Tamil, and Malayalam. Rajneesh Rani et al., in [27] reported a technique using Gabor filter and gradient based features using CNN classifier to recognize Gurumukhi and Roman scripts. The work was done at character level. From the literature review, it is observed that most of the work carried out is limited to only few Indian scripts and that the features employed either global (DCT, DWT, Gabor) or local (centroids, sphericity, aspect ratio, holes, connected component profiles, circularity). Further, local and global features have been used for both printed and handwritten text documents by Abdel Belad, Santosh K.C and Vincent Poulain d'Andecy in [28]. Local features are sensitive to noise that affect skewed correction and segmentation process. Global features fail to capture variations of the scripts that appear at smaller levels (subscripts or superscripts). Neural network based script recognition methodology is proposed by S Basavaraj Patil et al., in [29]. A two-stage feature extraction technique is incorporated. In the first stage, a 3x3 mask used to dilate image in all directions and at the second stage pixel distribution analysis is used. Experiments are performed to classify English, Hindi and Kannada scripts. Dhanya et al., in [30] have presented a two-stage approach to identify Roman and Tamil scripts present in bilingual documents. Character thickness, spatial spread in upper and lower zone of a word with three distinct spatial zones is used to identify script at the first stage and in the second stage respectively. A word level direction energy distribution using Gabor filters with appropriate frequency and orientations are used.

III. PROPOSED METHOD

The goal of handwriting recognition system is to process handwritten data electronically with the same or nearly the same accuracy as humans. By doing this process with computers, a large amount of data can be transcribed at high speed. An integrated approach to the design of OCR for all Indian scripts have great benefits. It is necessary to identify different script forms before running an individual OCR system. In a country like India, script identification is a must for multilingual OCR system. It acts as a pre-processor to the OCR system identifying the script type of the document, so that specific OCR tool can be selected as illustrated in Fig. 1. In a multi-script environment a bank of OCR corresponding to all different scripts are expected to be seen. The characters in an input document can then be recognized reliably by selecting the appropriate OCR system from the OCR bank.



Figure 1: Stages of document processing in a multi-script environment

A. Preprocessing

The pre-processed handwritten text document image in binary is considered as input to our segmentation algorithm. During the first stage, using horizontal profiles in, lines with little or no orientation of text-lines are extracted from the documents. The second stage corresponds to oriented lines (horizontal profiles of text-lines overlap) with certain degree of orientation. The third stage corresponds to those adjacent lines that are touching. Using a threshold value, overlapping adjacent text lines is detected and connected component technique is employed to extract these lines. The methodology is explained below. Horizontal profile of pre-processed document image is computed. The average height of text-line is determined based upon density of pixels in each line of the profile. Using this as a threshold, the document is parsed line by line to determine lines with no on pixels. The text-line lying between such lines are extracted subject to the condition that the height of the text-line meets the required threshold criteria of the average height of the text-line. In case, the height of the text-line exceeds the threshold, it is assumed that there are at least two overlapping lines, and such lines are subjected to second stage for text-line separation. By overlapping lines we mean, the on pixel density in the horizontal profile of neighborhood lines overlap and hence the height of the combined text-lines exceeds the average-height of the line making it difficult to separate the text-lines (Fig.2). Such lines usually arise due to the variation in writing text-lines by the writer.



Figure 2: Handwritten document image in Kannada with multi Document Image Segmentation

A. Feature Extraction

In this work the features like SIFT and Gabor are extracted from segmented scripts.

Scale Invariant Feature Transform (SIFT)

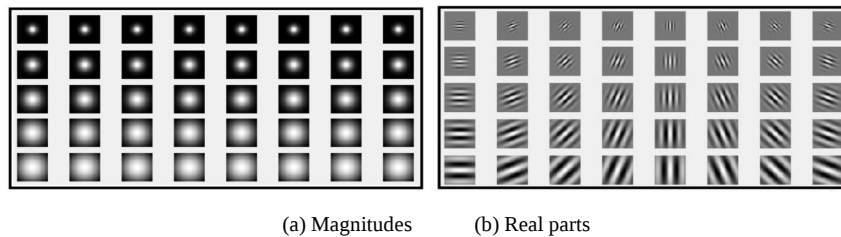
Scale-Invariant Feature Transformation (SIFT) is a widely used feature extraction technique in image classification problems. This feature is independent of the scale and orientation of the image and is robust to fluctuations in lighting, noise, partial overlap, and slight changes in image perspective. These properties are important for recognizing images of different sizes and orientations. The SIFT function consists of a number of image keypoints with orientation and corresponding area descriptors around the selected keypoint. SIFT key points are retrieved using a different image scale known as the Gaussian Difference Pyramid (DoG). The key point is selected by picking the point of the local maximum in the vicinity of the image scale pyramid. For each key point, a dominant direction is determined to achieve rotational invariance. SIFT descriptors are computed as image gradient histograms around keypoints to characterize the local occurrence of each selected keypoint. To help extraction of these Features, the SIFT calculation applies a 4 phase separating approaches:

1. Scale Spacing Extrema Detection
2. Localization of Key Point
3. Assigning Orientation
4. Key point Descriptor

Additionally, as the example above shows, close-by edge information can be used to create key point descriptors. Angle data were rotated to correspond with the direction of the key point and then a Gaussian with a $1.5 \times$ key point scale fluctuation was used to weight the data. Then, using a window that is focused on a critical point, this data is used to create numerous histograms. With 16 histograms displayed in a 4×4 frame, each with eight direction boxes, the key point descriptor typically employs one for each of the mid-purposes of these headings in addition to one for each of the basic compass bearings.

Gabor Feature Extraction Method

The Gabor wavelet features represent data at various frequencies and orientations suitable for derivation of discriminate features from logos. Gabor filters are directly related to Gabor wavelets and is a 2d kernel function. These features are generally used for perception of frequency with high components in the image leading to filtration of radiant details like textual component structures in the image. The filter is defined of real and imaginary components forming together the complex number representing various orthogonal directions. The details of Gabor filter bank is as depicted in figures 3 (a) and (b).



(a) Magnitudes (b) Real parts

Figure 3: Gabor filter bank

The convolution of character image with magnitudes and real parts of Gabor filter bank are the filtered images represented at 5 different scales and 8 different orientations and the same is as shown in figures 4(a) and (b).

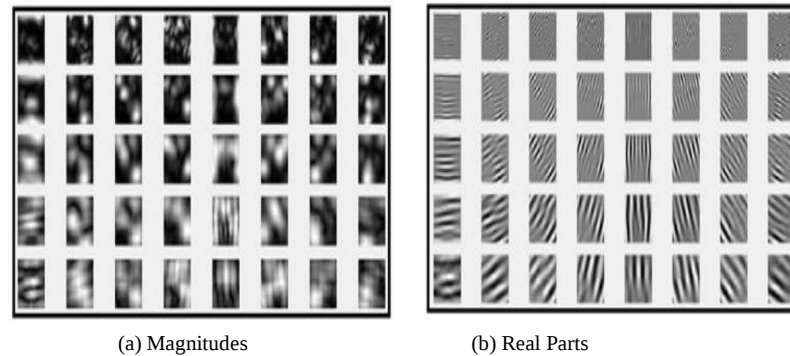


Figure 4: Filtered images

C. Classification

Classification of the script type is carried out using KNN and CNN classifiers. The KNN is a lazy classifier that compares the test image feature vector with that of feature vector of all the images used for training and labels the test image to be of a specific script type using Euclidean distance measure. The CNN classifier is a supervised learning method. Given a labeled training data (supervised learning), the algorithm outputs an optimal hyperplane which categorizes new examples.

K-Nearest Neighbor (KNN)

The K-Nearest Neighbors method is useful for solving classification and regression problems. The k closest training instances in the feature space make up the input in both scenarios. Typically, a distance algorithm based on data characteristics, such as the Euclidean distance formula, is used to calculate the proximity of two points. According to K values, the algorithm categorizes points based on the labels of their K nearest neighbours.

The KNN algorithm is performed by first determining the KNN parameter.

- a. Determine the separation between each training sample and each query instance.
- b. Use the minimal distance to K th to filter the distance and define the closest neighbours.
- c. Compile the category of your nearest neighbours.
- d. Using the predictive value as the simple majority of nearest neighbors group.

Convolutional Neural Network (CNN)

Deep learning is the neural network framework that can learn to infer knowledge from data. In deep learning, the reinforcement layer learns to convey input information into abstract representations as shown in Figure 5. Deep learning frameworks have advantages over traditional machine-learning methods such as computer vision and natural language processing because they are more flexible and allow for many nonlinear transformations. Convolutional neural networks (CNNs) are one of the proven DL frameworks in image classification and pattern recognition. CNNs were designed with inspiration from human brain activity, which is why their calculations are based on artificial neural networks. A CNN consists of a convolutional layer, a subsampling layer, and an output. The CNN architecture shown in Figure 4.4 uses a 32x32 pixel input image and has six layers. The first layer, called C1, employs six 5x5 convolutional kernels. Create 6 28x28 pixel functional maps from these kernels using a 2x2 pixel receptive field for each kernel. The second layer is the subsampling layer S2 with a 2x2 receiving field; create 16 feature maps from this layer using a 10x10 pixel size on each feature map. The third layer is the convolutional layer C3; it has sixteen 5x5 convolutional kernels. It generates sixteen feature maps each with a size of 10x10 pixels. The fourth layer is sub-sampling layer S4 with a 5x5 pixel size using a 2x2 pixel receptive field to create 16 feature maps from this layer using 10x10 pixels on each feature map. The fifth layer, the Convolutional Layer 5 (CL5), employs 120 5x5 kernel convolutions and generates 120 scalar features. A fully connected neural network has 120 input neurons and an output layer with a number of neurons based on the desired number of output classes. Initialization for each weight in CL5's kernels and layers is random. Back propagation is used to update filter weights during training, which is more semantically rich than raw images.

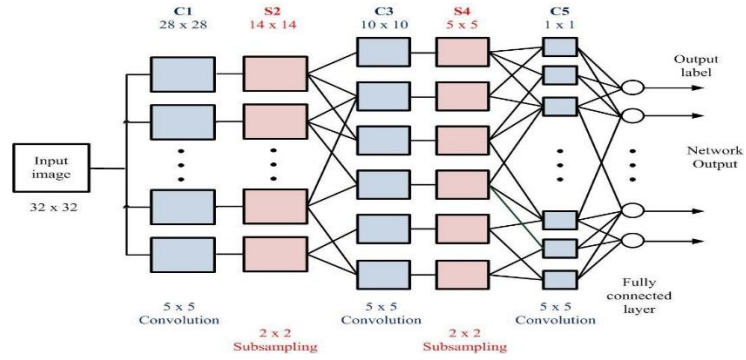


Figure 5: Architecture of CNN used for classification

IV. EXPERIMENTAL RESULTS

Experiments are carried out on scripts types viz, Kannada, Hindi, Urdu, Malayalam, Punjabi, Telugu, Tamil, Oriya and English respectively. A total of 230 handwritten documents are considered for experimentation. Text lines with only 40% of occupancy of text are rejected by considering it as not a complete line. Sample images of text lines rejected are shown in Fig 5, 630 text line images are used to test the proposed method. The overall recognition accuracy obtained using K-NN classifier is 92.63% and 94.9% accuracy is obtained using CNN classifier. Two-fold cross-validation is performed for computing the accuracy of the classification results. The details of the results obtained for bi-script documents are presented in Table 1. The results obtained for tri-script classification is shown in Table 2. CNN classifier performs better over KNN classifier.

TABLE I : BI-SCRIPT CLASSIFICATION USING KNN AND CNN CLASSIFIERS

Sl. No.	Script	KNN	Recognitionin %	CNN	Recognitionin %
1	Kannada	92.86	95.71	95.71	97.85
	English	98.57		100	
2	Hindi	98.57	98.57	100	99.28
	English	98.57		98.57	
3	Malayalam	92.86	93.57	95.71	90.71
	English	94.29		85.71	
4	Oriya	92.86	90	98.57	95
	English	87.14		91.43	
5	Punjabi	90	90	95.71	92.85
	English	90		90	
6	Tamil	94.29	91.43	95.71	90.71
	English	88.57		85.71	
7	Telugu	95.71	94.96	100	97.85
	English	94.21		95.71	
8	Urdu	91.43	90.71	100	93.57
	English	90		87.14	
Over all Recognitionaccuracy in %			93.12		94.73

TABLE II: TRI-SCRIPT CLASSIFICATION USING KNN AND CNN CLASSIFIERS

Sl. No.	Script	KNN	Recognitionin %	CNN	Recognitionin %
1	Kannada	95.71	96.66	92.86	97.14
	Hindi	98.57		100	
	English	95.71		98.57	
2	Malayalam	100	92.86	100	94.76
	Hindi	84.29		94.29	
	English	94.29		90	
	Oriya	95.71		97.14	

3	Hindi	88.57	90.47	94.29	94.29
	English	87.14		91.43	
4	Tamil	98.57	87.62	100	90.01
	Hindi	91.43		88.59	
	English	72.86		81.43	
5	Telugu	97.14	93.81	100	98.57
	Hindi	91.43		100	
	English	92.86		95.71	
6	Urdu	97.14	91.43	100	95.71
	Hindi	88.57		98.57	
	English	88.57		88.57	
Over all Recognition accuracy in %			92.14		95.08

Handwritten Script Identification at Word level

Experiments are carried out on scripts types viz, Kannada, Hindi, Urdu, Malayalam, Telugu, Tamil, Oriya and English respectively. A total of 2400 handwritten text word images i.e. Two characters word, Three characters word and More than Three characters word (800 word images each) are considered for experimentation. The overall recognition accuracy obtained using K-NN classifier is 92.28% and 95.77% accuracy is obtained using CNN classifier. Two-fold cross-validation is performed for computing the accuracy of the classification results. The details of the results obtained for bi-script documents are presented in Table 4. The results obtained for tri-script classification is shown in Table 6. CNN classifier performs better over KNN classifier and the results are shown in table 5.

TABLE III: RECOGNITION RESULT IN %(BI-SCRIPT)

Script	Bi-Script					
	Two Character Word		Three Character Word		More Three Character Word	
	KNN	CNN	KNN	CNN	KNN	CNN
English, Hindi	100	100	100	100	98	100
English, Kannada	100	100	99.5	99.5	92.5	97
English, Malayalam	70	79.5	82	87	86.5	94.5
English, Oriya	100	100	100	100	83.5	86.5
English, Tamil	89.5	93.5	100	100	76	88
English, Telugu	80	91	100	100	97	99.5
English, Urdu	94.5	95	100	100	96	98
Recognition accuracy in %	90.57	94.14	97.36	98.07	89.93	94.79

TABLE IV: RECOGNITION RESULT IN %(TRI-SCRIPT)

Script	Tri-Script					
	Two Character Work		Three Character Word		More Three Character Word	
	KNN	CNN	KNN	CNN	KNN	CNN
English, Hindi, Kannada	89	98	97.3	99	90.3	99
English, Hindi, Malayalam	81	88.3	88.3	92.7	88.7	97.3
English, Hindi, Oriya	91.7	96.7	97	98.7	89	92
English, Hindi, Tamil	91.7	91.7	97.7	98.7	82.7	90.7
English, Hindi, Telugu	90.7	92.3	98.7	99.3	99.3	99.7
English, Hindi, Urdu	96.7	95.3	99.7	99.3	85.3	97
Recognition accuracy in %	90.13	93.72	96.45	97.95	89.22	95.95

TABLE V: OVER ALL RECOGNITION FOR BI-SCRIPT AND TRI-SCRIPT

Over all accuracy for Bi-script and Tri-script		
	KNN	CNN
Bi-Script	92.62	95.67
Tri-Script	91.93	95.87
Over all accuracy in %	92.28	95.77

V. CONCLUSION

Script type identification using Gabor and SIFT, a powerful feature for texture classification, has been presented in this chapter. Experiments are performed on scriptimages at line, word, and block level. Classification results are computed using KNN and CNN classifiers at bi-script and tri-script level. For images of text line, K-NN yielded overall accuracy of 92.63%, and CNN yielded over all accuracy of 94.91%. For word level images, K-NN yielded overall accuracy of 92.28%, and CNN yielded over all accuracy of 95.77%. Finally, for block level images, of 512x512 pixels having 3-lines, 4-lines and more than 4-lines of text, K-NN yielded 98.46 % accuracy and 99.5% accuracy was observed with CNN. It is observed that performance ofCNN classifier is better over K-NN classifier in terms of overall recognition of script type identification in all the image types.

REFERENCES

- [1] B.B. Chaudhuri, U. Pal, A Complete Printed Bangla OCR System, Pattern Recognition. 31 (1998) 531–549.
- [2] U. Pal, B.B. Chaudhuri, Indian Script Character Recognition: A survey, Pattern Recognition. 37 (9) (2004) 1887–1899.
- [3] Dan Claudiu Cires, An And Ueli Meier And Luca Maria Gambardella And Jurgen Schmid Huber, “Convolutional Neural Network Committees For Handwritten Character Classification”, 2011 International Conference On Document Analysis And Recognition, IEEE, 2011.
- [4] Shrey Dutta, Naveen Sankaran, Pramod Sankar K., C.V. Jawahar, “Robust Recognition Of Degraded Documents Using Character N-Grams”, IEEE, 2012.
- [5] Sukhpreet Singh " Optical Character Recognition Techniques: A Survey", ISSN: 2278 – 1323 International Journal Of Advanced Research In Computer Engineering & Technology (IJARCET) Volume 2, Issue 6, June 2013.Pp 2009-2015.
- [6] Pawan Kumar Singh, Ram Sarkar, Mita Nasipuri, Offline Script Identification From Multilingual Indic-Script Documents: A State-Of-The-Art ,Computer Science Review, Volumes 15–16, February – May 2015, Pages 1-28.
- [7] Sahare, P., & Dhok, S. B. (2017). Script Identification Algorithms: A Survey. International Journal Of Multimedia Information Retrieval, 6(3), 211–232. Doi:10.1007/S 13735-017-0130-2.
- [8] K. Ubul, G. Tursun, A. Aysa, D. Impedovo, G. Pirlo, T. Yibulayin , “ Script Identification Of Multi-Script Documents: A Survey ” , IEEE Access DOI 10.1109/ACCESS.2017.2689159 Volume 5 2017, Pp. 6546-6559.
- [9] S. Ben Moussa, A. Zahour, A. Benabdelhafid And A.M. Alimi, “Fractal-Based System For Arabic/Latin, Printed/Handwritten Script Identification, ” 978-1-4244-2175- 6/08/\$25.00 ©2008 IEEE.
- [10] M. C. PADMA, P. A. VIJAYA “SCRIPT IDENTIFICATION FROM TRILINGUAL DOCUMENTS USING PROFILE BASED FEATURES ” International Journal Of Computer Science And Applications, Technomathematics Research Foundation Vol. 7 No. 4, Pp. 16 - 33 , 2010.
- [11] Mallikarjun Hangarge And B.V.Dhendra “Offline Handwritten Script Identification In Document Images” , International Journal Of Computer Applications (0975 – 8887) Volume 4 – No.6, July 2010. 131
- [12] Dr. G.G. Rajput, Anita H.B “Handwritten Script Recognition At Line Level – A Multiple Feature Based Approach” International Journal Of Engineering And Innovative Technology (IJEIT) Volume 3, Issue 4, October 2013. Pp 90-95.
- [13] Prasanthkumar P V And Dileesh E D “Word Level Script And Language Identification For Unconstrained Handwritten Document Images” 2014 3rd International Conference On Eco-Friendly Computing And Communication Systems , 978-1-4799-7002-5/14 \$31.00 © 2014 IEEE DOI 10.1109/Eco-Friendly.2014.78,
- [14] Adnan Ul-Hasan, Muhammad Zeshan Afzal, Faisal Shafait, Marcus Liwicki, And Thomas M. Breuel “A Sequence Learning Approach For Multiple Script Identification”. 978-1-4799-1805-8/15/\$31.00 ©2015 IEEE Pp 1046-1050.
- [15] Sk Md Obaidullah, Rownaqul Karim, Sujal Shaikh, Chayan Halder, Nibaran Das, Kaushikroy “Transform Based Approach For Indic Script Identification From Handwritten Document Images” 2015 3rd International Conference On Signal Processing, Communication And Networking (ICSCN), 978-1-4673-6823-0/15/\$31.00 ©2015 IEEE.
- [16] Sk Md Obaidullah, Nibaran Das, Kaushik Roy “ Convolution Based Technique For Indic Script Identification From Handwritten Document Images” IJ. Image, Graphics And Signal Processing, 2015, 5, 49-57 Published Online April 2015 In 12. MECS (Http://Www.Mecs-Press.Org/) DOI: 10.5815/Ijigsp.2015.05.06.
- [17] Baoguang Shi. Xiang Bai , Cong Yao “Script Identification In The Wild Via Discriminative Convolutional Neural Network” Pattern Recognition 52 (2016)448–458.

- [18] Sk Md Obaidullah , K. C. Santosh, Chayan Halder, Nibaran Das ,Kaushik Roy “Automatic Indic Script Identification From Handwritten Documents: Page, Block, Line And Word Level Approach” Int. J. Mach. Learn. & Cyber. DOI 10.1007/S13042-017- 0702-8. Published Online 2017.
- [19] Sk Md Obaidullah, Chayan Halder2, K. C. Santosh, Nibaran Das, Kaushik Roy “Automatic Line-Level Script Identification From Handwritten Document Images - A Region-Wise Classification Framework For Indian Subcontinent” Malaysian Journal Of Computer Science. Vol. 31(1), 2018 , Pp 63-84.
- [20] Judith Hochberg, Kevin Bowers, Michael Cannon And Patrick Kelly, “Script And Language Identi_Cation For Handwritten Document Images,” International Journal On Document Analysis And Recognition ©Springer-Verlag 1999. IJDAR (1999) 2: 45- 52. 132
- [21] V. Singhal, N. Navin And D. Ghosh , “Script-Based Classification Of Hand-Written Text Documents In A Multilingual Environment,” 0-7803-7868-7/03/\$17.00 ©2003 IEEE.
- [22] DAVID G. LOWE, “Distinctive Image Features From Scale-Invariant Keypoints,” International Journal Of Computer Vision 60(2), 91–110, 2004 © 2004 Kluwer Academic Publishers. Manufactured In The Netherlands.
- [23] Lijun Zhou, Yue Lu And Chew Lim Tan , “Bangla/English Script Identification Based On Analysis Of Connected Component Profiles,” H. Bunke And A.L. Spitz (Eds.): DAS 2006, LNCS 3872, Pp. 243–254, 2006. ©Springer-Verlag Berlin Heidelberg 2006.
- [24] G. G. Rajput And Anita H. B , “Handwritten Script Recognition Using DCT And Wavelet Features At Block Level,” IJCA Special Issue On “Recent Trends In Image Processing And Pattern Recognition RTIPPR, 2010.Pp 158-163.
- [25] K. Roy, S. Kundu Das And Sk Md Obaidullah, “Script Identification From Handwritten Document,” 2011 Third National Conference On Computer Vision, Pattern Recognition, Image Processing And Graphics, 978-0-7695-4599-8/11 \$26.00 © 2011 IEEE DOI 10.1109/NCVPRIPG.2011.22.
- [26] Mallikarjun Hangarge, K.C. Santosh And Rajmohan Pardeshi,” Directional Discrete Cosine Transform For Handwritten Script Identification,” 2013 12th International Conference On Document Analysis And Recognition , 1520-5363/13 \$26.00 © 2013 IEEE ,DOI 10.1109/ICDAR.2013.76.
- [27] Rajneesh Rani, Renu Dhir And Gurpreet Singh Lehal , “Script Identification Of Presegmented Multi-Font Characters And Digits,” 2013 12th International Conference On Document Analysis And Recognition, 1520-5363/13 \$26.00 © 2013 IEEE ,DOI 10.1109/ICDAR.2013.233.
- [28] Abdel Belad, Santosh K.C And Vincent Poulain d'Andecy, “Handwritten And Printed Text Separation In Real Document,” The Thirteenth IAPR International Conference On Machine Vision Applications - 2013, May 2013, Kyoto, Japan. 2013.
- [29] S Basavaraj Patil, N. V. Subbareddy, "Neural Network-Based System For Script Identification In Indian Documents", Sadhana Vol. 27, Part 1, February 2002, Pp. 83– 97.
- [30] D Dhanya, A G Ramakrishnan_ And Peeta Basa Pati “Script Identification In Printed Bilingual Documents” , Sadhana Vol. 27, Part 1, February 2002, Pp. 73–82.