

Handwritten English Character Recognition and translate English to Devnagari Words

Shivali Parkhedkar, Shaverai Vairagade, Vishakha Sakharkar, Bharti Khurpe, Arpita Pikalmunde, Amit Meshram, Prof. Rakesh Jambhulkar

Department of Information Technology , RTMNU/MIET Bhandara, Maharashtra, India

ABSTRACT

In our proposed work we will accept the challenges of recognizing the words and we will work to win the challenge. The handwritten document is scanned using a scanner. The image of the scanned document is processed victimization the program. Each character in the word is isolated. Then the individual isolated character is subjected to “Feature Extraction” by the Gabor Feature. Extracted features are passed through KNN classifier. Finally we get the Recognized word. Character recognition is a process by which computer recognizes handwritten characters and turns them into a format which a user can understand. Computer primarily based pattern recognition may be a method that involves many sub process. In today’s surroundings character recognition has gained ton of concentration with in the field of pattern recognition. Handwritten character recognition is beneficial in cheque process in banks, form processing systems and many more. Character recognition is one in all the favored and difficult space in analysis. In future, character recognition creates paperless environment. The novelty of this approach is to achieve better accuracy, reduced computational time for recognition of handwritten characters. The proposed method extracts the geometric features of the character contour. These features are based on the basic line types that forms the character skeleton. The system offers a feature vector as its output. The feature vectors so generated from a training set, were then used to train a pattern recognition engine based on Neural Networks so that the system can be benchmarked. The algorithm proposed concentrates on the same. It extracts totally different line varieties that forms a specific character. It conjointly also concentrates on the point options of constant. The feature extraction technique explained was tested using a Neural Network which was trained with the feature vectors obtained from the proposed method.

Keywords : HCR(Handwritten Character Recognition), QBT(Querriy By Text), QBS(Querriy By String), DTW(Dynamic Time Wrapping), ICA(Independent Component Analysis).

I. INTRODUCTION

TEXT that has drawn a lot of attention from the computer understanding in images is an important problem vision community since its beginnings. Text understanding covers several applications and tasks, most of that originated decades ago thanks to the digitalization of huge collections of documents. This create necessary the event of ways ready to extract

info from these document images: layout analysis, information flow, transcription and localization of words, etc. Recently, and intended by the exponential increase of publically out there image databases and private collections of images, this interest now also embraces text understanding on natural images. Methods ready to retrieve pictures containing a given word or to acknowledge words during a image have additionally become possible and

helpful. In this paper we have tendency to take into account 2 issue associated with text understanding: word recognizing and word recognition. In word recognizing, the goal is to find all instances of a query word in a dataset of images. The question word is also a text string—in that case it's sometime noted as question by string (QBS) or question by text (QBT)—, or may also be an image, during which case it's sometimes noted as question by example (QBE). In word recognition, the goal is to obtain a transcription of the query word image. In several cases, together with this work, it's assumed that a text wordbook or lexicon is provided at take a look at time, which solely words from that lexicon will be used as candidate transcriptions in the recognition task. In this work

we will additionally assume that the situation of the words with in the pictures is provided, i.e., we've access to photographs of cropped words. If those weren't on the market, text localization and segmentation techniques might be used. historically, word recognizing and recognition have centred on document pictures, where the most challenges return from differences in writing styles: the writing varieties of totally different{completely different}writers is also fully different for a similar word. Recently, however, with the event of powerful laptop vision techniques throughout the last decade, there has been associated degree multiplied interest in performing arts word recognizing and recognition on neural pictures which poses different challenges such as huge variations in illumination, point of view, typography, etc. Word spotting can be seen as a particular case of semantic content based image retrieval (CBIR), where the classes are very fine-grained, we are interested in exactly one particular word, and a distinction of just one character is taken into account a negative result—but additionally contain a awfully giant intra-class variability, writing designs, illumination, typography, etc., can make the

same word look very different. In the same means, word recognition will be seen as a special case of terribly fine-grained, zero-shot classification, where we are interested in classifying a word image into (potentially) hundreds of thousands of categories, that we tend to might not have seen any coaching example. By having pictures and text strings share a standard mathematical space with an outlined metric, word recognizing and recognition become an easy nearest neighbour drawback during a low dimensional house. We can perform QBE and QBS (or even a hybrid QBE, where both an image and its text label are provided as queries) using exactly the same retrieval framework. The recognition task simply becomes finding the nearest neighbour of the image word in a text dictionary embedded first into the PHOC space and then into the common subspace. Since we use compact vectors, compression and indexing techniques such as Product Quantization could now be used to perform spotting in very large datasets. To the best of our knowledge, we are the first to provide a unified framework where we can perform out of vocabulary QBE and QBS retrieval as well as word recognition using the same compact word representations. We review the related work in word spotting and recognition. Word recognizing in document pictures has attracted attention with in the document analysis community throughout the last 20 years and still poses countless challenges thanks to the difficulties of historical documents, different scripts, noise, handwritten documents, etc. Because of this complexness, most popular techniques on document word spotting have been based on describing word images as sequences of features of variable length and using techniques such as dynamic time warping (DTW) or Hidden mathematician(HMM) to classify them. Variable length features are more flexible than feature vectors and have been known to lead to superior results in difficult word-spotting tasks since they can adapt better to the different variations of style and word length. The increasing interest in extracting matter

info from real scenes is said to the recent growth of image databases like Google pictures or Flickr. Some attention-grabbing tasks are recently planned, e.g., localization and recognition of text in Google Street read pictures or recognition in signs harvested from Google pictures. The high complexness of those pictures when put next to documents, mainly due the large appearance variability, makes it very difficult to apply traditional techniques of the document analysis field. However, with the recent development of powerful laptop vision techniques some new approaches are planned. Some methods have focused on the problem of end-to-end word recognition, which comprises the tasks of text localization and recognition. They first detect a set of possible character candidates windows using a sliding window approach and then each word in the lexicon is matched to these detections. Finally, the one with the best score is rumored because the foretold word. Then, the recognition of candidate regions is done in a separate OCR stage using synthetic fonts. They first perform a text detection process returning candidate regions containing individual lines of text, which are then processed for text recognition. This recognition is done by identifying candidate character regions labels is defined. Of those, since it also deals with text recognition and presents a text embedding approach (spatial pyramid of characters or SPOC) very similar to ours. The main distinction stems from however the embedding is employed. While in our case, we use it as a source of attributes, and only then we try to find a common subspace between the attributes. Our approach will be seen as a additional regular version of theirs, since we enforce that the projection that embeds our images into the common subspace can be decomposed into a matrix that comes the photographs into associated degree attributes house. This work present associated degree Offline Cursive Word Recognition System managing single author samples. The system relies on both technique significantly improve the popularity rate of the system. Pre-processing, normalization and feature

extraction are described as well as the training technique adopted. Several experiments were performed employing publically on the market information. The accuracy obtained is the highest presented in the literature over the same data. The system is based on a sliding window approach: a window shifts column by column across the image and, at each step, isolates a frame. A feature vector is extracted from every frame and therefore the sequence of frames thus obtained is sculptured with Continuous Density Hidden Mathematician model (HMMs). The use of the sliding window approach has the important advantage of avoiding the need of an independent segmentation, a difficult and error prone process. In order to reduce the number of parameters in the HMMs, we use diagonal covariance matrices in the emission probabilities. This corresponds to the unrealistic assumption of having de co-related feature vectors. For this reason, we applied Principal Component Analysis (PCA) and Independent Component Analysis (ICA) to de-correlate the data. This allowed a significant improvement of the recognition rate. The recognition accuracy achieved with the approach proposed here is to our knowledge, the highest among the results over the same data presented in the literature. The analysis of the recognition as a function of the word length shows that the system achieves a recognition rate for samples longer than six letters. This suggests that the performance of our system in tasks involving words with high average length are often superb. Both PCA and ICA had a positive result on the recognition rate, PCA in particular reduced the error rate. A further improvement will most likely be obtained by victimization nonlinear or kernel PCA. Such techniques typically work higher than the linear remodel we tend to accustomed PCA. The use of data dependent heuristics was avoided in order to make the system flexible with respect to a change of writer. Any ad-hoc algorithm for the specific style of the writer was avoided. The prior information about the word frequency and distribution can be useful to

improve the recognition of short words. These are typically articles, conjunctions and propositions that appear often in the sentences. For this reason, a possible future direction to follow is the application of language models that take into account this kind of information. Point Times New Roman for Figure and Table labels. Use words instead of symbols or abbreviations once writing Figure axis labels to avoid confusing the reader.

II. METHODS AND MATERIAL

A. EXISTING SYSTEM

This work presents an Offline Cursive Word Recognition System addressing single author samples.[58] The system is predicated on a both techniques a continual improved the popularity rate of the system. density Pre-processing, normalization and feature extraction are described as well as the training technique adopted. Several experiments were performed employing a in public accessible information. The accuracy obtained is that the highest given within the literature over constant information. The system is based on a sliding window approach: a window shifts column by column across the image and, at each step, isolates a frame. A feature vector is extracted from each frame and the sequence of frames so obtained is modelled with Continuous Density Hidden Markov Models (HMMs). The use of the sliding window approach has the important advantage of avoiding the need of an independent segmentation, a difficult and error prone process. In order to reduce the number of parameters in the HMMs, we use diagonal covariance matrices in the emission probabilities. This corresponds to the unrealistic assumption of having de-correlated feature vectors. For this reason, we applied Principal Component Analysis (PCA) and Independent Component Analysis (ICA) to de-correlate the data. This allowed a significant improvement of the recognition rate. The recognition accuracy achieved

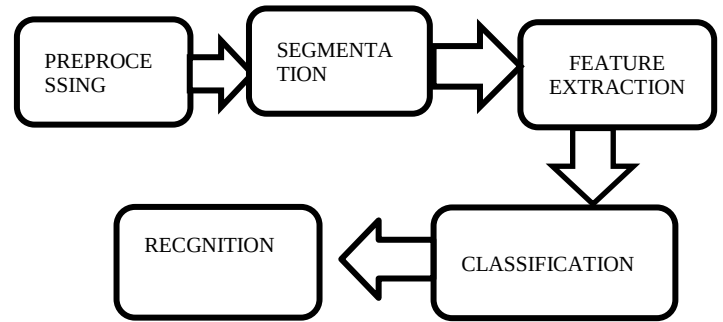
with the approach proposed here is to our knowledge, the highest among the results over the same data presented in the literature. The analysis of the recognition as a function of the word length shows that the system achieves a recognition rate for samples longer than six letters. This suggests that the performance of our system in tasks involving words with high average length is excellent. Both PCA and ICA had a positive impact on the popularity rate, PCA in particular reduced the error rate. A further improvement will in all probability be obtained by mistreatment nonlinear or kernel PCA. Such techniques typically work higher than the linear rework we tend to want to perform PCA. The use of data dependent heuristics was avoided in order to make the system flexible with respect to a change of writer. Any ad-hoc algorithm for the specific style of the writer was avoided. The prior information about the word frequency and distribution can be useful to improve the recognition of short words. These are typically articles, conjunctions and propositions that appear often in the sentences. For this reason, a possible future direction to follow is the application of language models that take into account this kind of information.

B. PROPOSED SYSTEM

In the proposed system, this is achieved by a combination of label embedding and attributes learning, and a common subspace regression. Then the images and strings represent the same word which are close to each other allowing one to cast recognition and retrieval tasks. Compared with the existing method, the advantage of our method has a fixed length, low dimensional and very fast to compute. Word spotting in document images has attracted attention in the document analysis and still poses lots of challenges due to the difficulties of historical documents, different scripts, noise, handwritten documents, etc. Regarding word

recognition, handwritten recognition still poses an important challenge for the same reasons. A model is first trained using labelled training data. At test time, given an image word and a text word, the model computes the probability of that text word being produced by the model when fed with the image word. Recognition will then be self-addressed by the chance of all the lexicon words given the question image and retrieving the closest neighbour. As in the word spotting case, the main drawback here is the comparison speed, since computing these probabilities is orders of magnitude slower than computing a Euclidean distance or a dot product between vectorial representations. The increasing interest in extracting matter data from real scenes is expounded to the recent growth of image databases like Google pictures or Flickr. Some attention-grabbing tasks are recently planned, e.g., localization and recognition of text in Google Street View images or recognition in signs harvested from Google Images. The high complexness of those pictures when put next to documents, mainly due the large appearance variability, makes it very difficult to apply traditional techniques of the document analysis field. However, with the recent development of powerful laptop vision techniques some new approaches are planned. To learn a way to retrieve and recognize words that haven't been seen during training, it's necessary to be ready to transfer data between the coaching and testing samples. One of the most popular approaches to perform this zero shot learning in computer vision involves the use of visual attributes in our case, we propose a broader framework since we do not constrain the choice of features or the method to learn the attributes.

SYSTEM ARCHITECTURE



MODULE

- Pre-processing
- Segmentation
- Feature Extraction (Gabor)
- Classification and recognition

C. MODULE DESCRIPTION

PREPROCESSING

First, the image is loaded because the input image. Filtering operations takes place for the input image. The median filter is used for the removal of noise and smoothening the image. The median filter could be a nonlinear digital filtering technique, typically wont to take away noise. Such noise reduction could be a typical pre-processing step to enhance the results of later process. (for example, edge detection on associated degree image). Median filtering is extremely wide utilized in digital image process as a result of under certain conditions. it preserves edges while removing noise. Median filtering smooth the image and is thus useful in reducing noise. Unlike low pass filtering, median filtering can preserve discontinuities in a step function and can smooth a few pixels whose values differ significantly from their surroundings without affecting the other pixels. It is conjointly helpful in conserving edges in a picture where as reducing random noise. Impulsive or salt-and pepper noise will occur thanks to a random bit error in an exceedingly communicating.

SEGMENTATION

We take into account 2 issue associated with text understanding: word recognizing and word recognition. In word recognizing, the goal is to search out all instances of a question word in a very dataset of pictures. The question word is also a text string—in that case it's typically stated as query by string (QBS) or query by text (QBT)—,or may also be an image,—in that case it's typically stated as query by example (QBE). In word recognition, the goal is to get a transcription of the query word image. In several cases, together with this work, it's assumed that a text dictionary or lexicon is provided at take a look at time, which solely word from that lexicon are often used as candidate transcriptions in the recognition task. In this work we'll conjointly assume that the placement of the words within the images is provided, i.e., we've access to images of cropped words. In this work we'll conjointly assume that the placement of the words within the picture is provided, i.e. we've access to photographs of cropped words. In the segmentation process we are cropping the words identically and show it in the bounding box

FEATURE EXTRACTION

Gabor wavelets in image process algorithms, particularly the interest purpose detection. There are many approaches to the interest purpose detection victimization Gabor functions or wavelets. More specifically, the two most common approaches involve the edge detection from the feature image or the corner detection using a combination of responses to several filters with a different orientation . In this paper, a new method based on the Gabor wavelets is proposed. This approach differs from previous approaches primarily within the manner the filter response is computed. More specifically, the filter response is decided solely in two perpendicular directions. The nature of this

approach within the use of responses to Gabor wavelets because the partial derivatives within the partial detectors. (e.g., Canny edge detector, Harris corner detector, Hessian-based blob detector). Such an approach may be useful when a fast implementation of the Gabor transform is available, or when the transform is already pre-computed. My main contribution consists of the use of the Gabor wave as a multi scale partial differential operator.

CLASSIFICATION

The classification for this process is done by using KNN Classifier. Classification (generalization) victimization associate in nursing instance-based classifier are often an easy matter of locating the closest neighbour in instance house(space) and labelling the unknown instance with the identical class label as that of the located (known) neighbour. This approach is commonly mentioned referred to as a nearest neighbour classifier. The drawback of this straightforward approach is that the lack of lustiness that characterizes the ensuing classifiers. The high degree of native sensitivity makes nearest neighbour classifiers extremely liable to noise within the coaching information.

KNN CLASSIFIER

In word recognition, the k-nearest neighbour algorithm (k-NN) is a non-parametric method for classifying objects based on closest training examples in the feature space. In k-NN the function is only approximated locally and all computation is deferred until classification. The k-nearest neighbour algorithmic program is amongst the only of the machine learning algorithms: An object is accessed by a majority vote of its neighbours, with the object being assigned to the category commonest amongst its k nearest neighbours (k may be a positive integer, typically small). If $k = 1$, then the article is solely allotted to the category of that single nearest

neighbour. The same methodology may be used for regression, by simply assigning the property value for the object to be the average of the values of its k nearest neighbours. It can be helpful to weight the contributions of the neighbours, so that the nearer neighbours contribute more to the average than the more distant ones. (A common coefficient scheme is to offer every neighbour a weight of $1/d$, where d is the distance to the neighbour. This scheme is a generalization of linear interpolation.). The neighbours are taken from a group of objects that the proper classification (or, in the case of regression, the value of the property) is known.

The learning process

Unlike many artificial learners, instance-based learners do not abstract any information from the training data during the learning phase. Learning is simply a matter question of encapsulating the training data. The process of generalization is delayed till it's completely unavoidable, that is, at the time of classification. This property has result in the concerning instance-based learners as lazy learners, whereas classifiers like feed forward neural networks, whenever correct abstraction is completed throughout the training phase often are entitled eager learners.

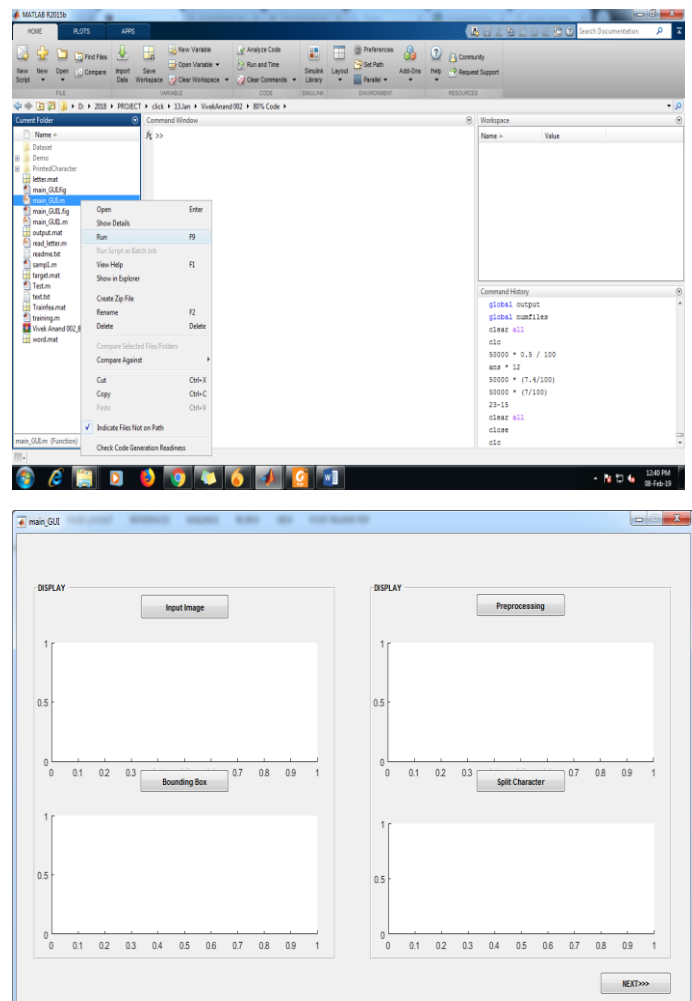
ADVANTAGE:

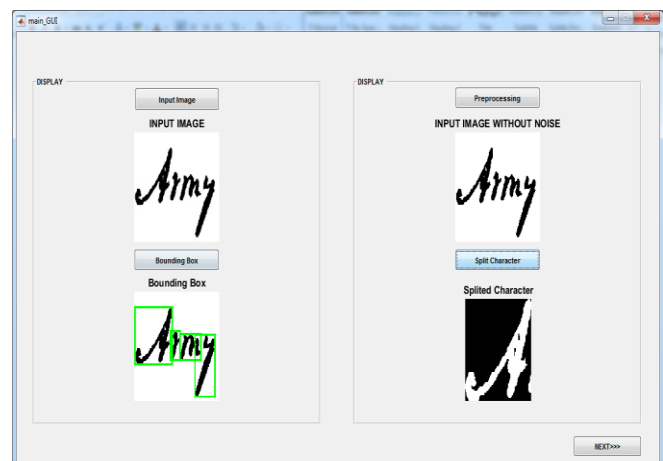
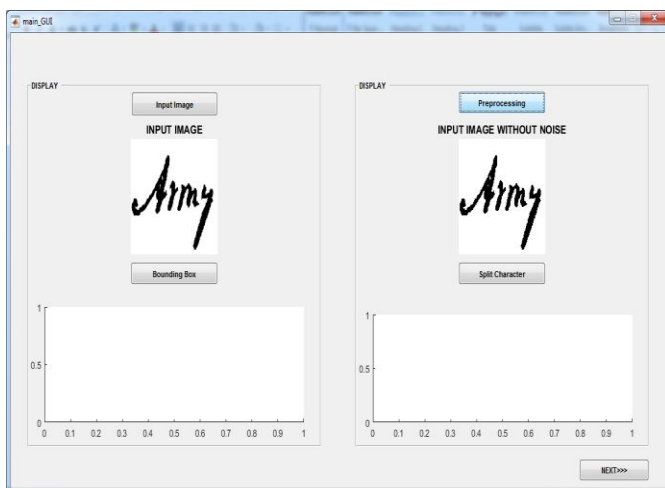
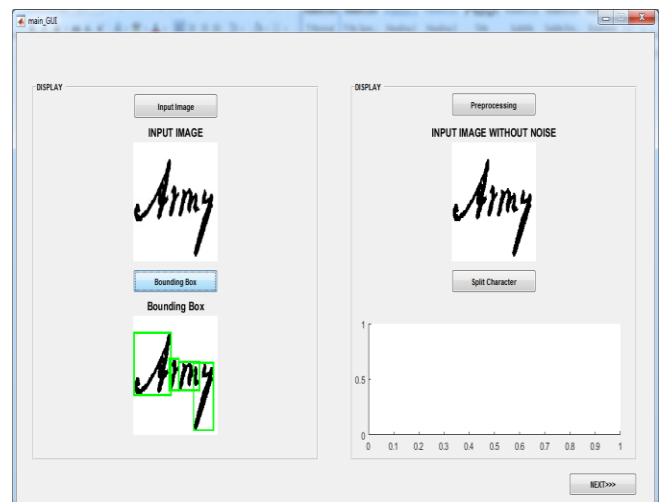
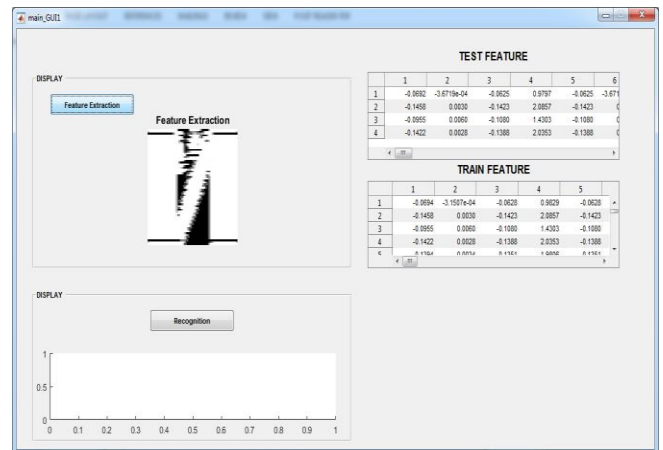
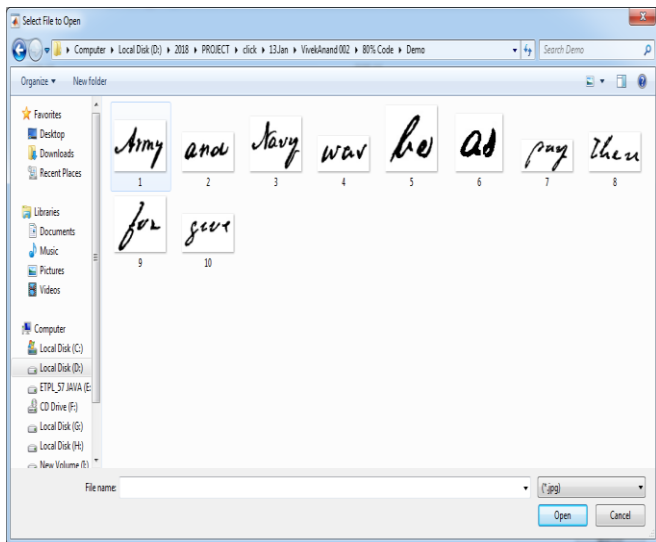
- Our learning approach is currently based on whole word images and does not require to segment the individual characters of the words during training or test, and leads to the large improvements in accuracy.
- This method does not require an available lexicon for a full recognition of the image words.

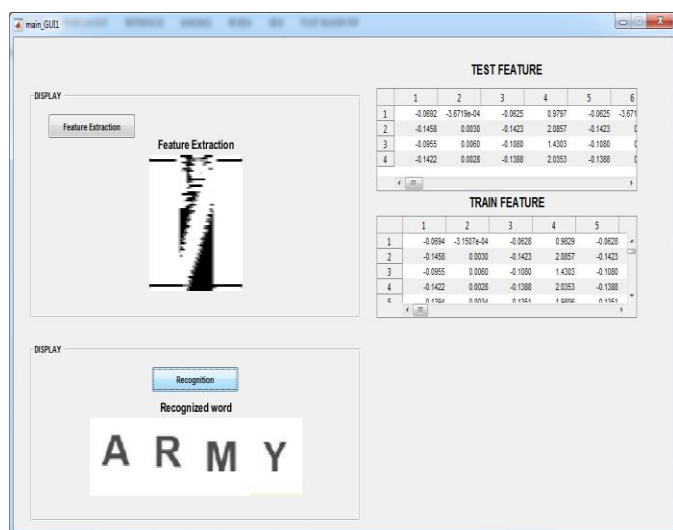
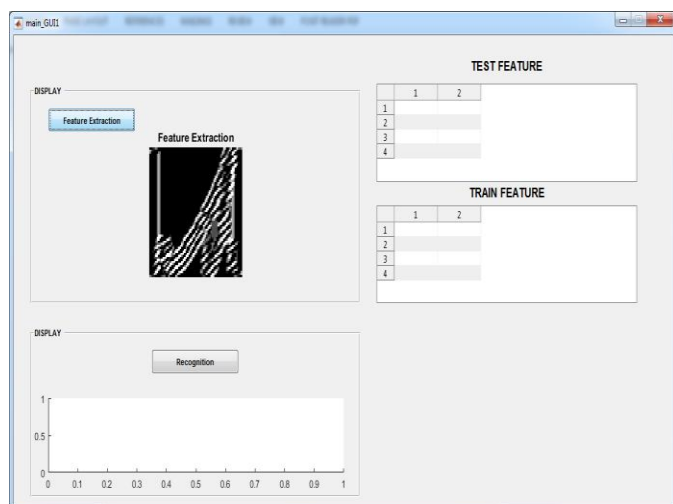
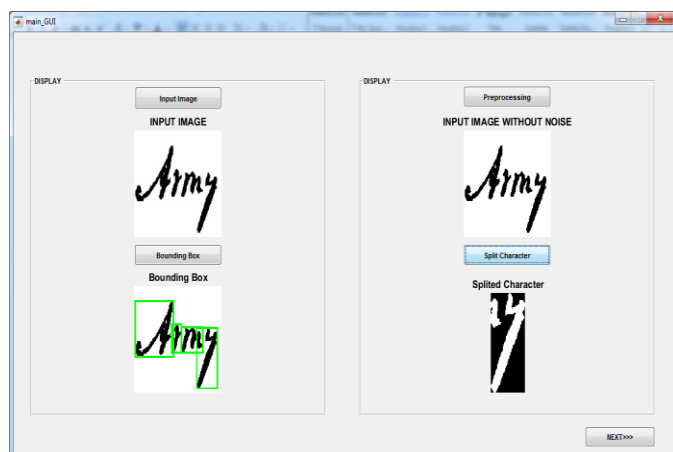
DISADVANTAGE:

- The problem is that the standard neural network objective functions are defined separately for each point in the training sequence .

III. RESULTS







IV. CONCLUSION

This paper proposes an approach to represent and compare word images, both on document and on natural domains. We show however an attributes-based approach supported a pointed bar chart of

characters may be wont to find out how to implement the word picture and their matter transcription into a shared, additional discriminative space, whenever the similarity between words is freelance of the writing and font vogue, illumination, capture angle, etc. This attributes representation leads to a unified representation of word images and strings, resulting in a method that allows one to perform either query-by-example or query-by-string searches, as well as image transcription, in a unified framework. We test our method in four public datasets of documents and natural images, outperforming state-of-the-art approaches and showing that the proposed attribute-based representation is well suited for word searches, whether they are pictures or strings, in handwritten and natural images. The results of our approach on the word spotting task .we compare the FV baseline (which can only be used in QBE tasks), the uncalibrated attributes embedding (Att.), the attributes calibrated with Plates (Att. + Platts), the one-way regression (alt.+Reg), the common subspace regression (Att. + CSR), and the kernelized common subspace regression (Att. + KCSR).

[1]. REFERENCES

- [2]. J Almazan, A. Gordo, A. Fornes, and E. Valveny, "Handwritten word spotting with corrected attributes," in Proc. IEEE Int. Conf. Comput. Vis., 2013, pp. 1017–1024.
- [3]. R Manmatha and J. Rothfeder, "A scale space approach for automatically segmenting words from historical handwritten documents," IEEE Trans. Pattern Anal. Mach. Intell., vol. 27, no. 8, pp. 1212–1225, Aug. 2005.
- [4]. L Neumann and J. Matas, "Real-time scene text localization and recognition," in Proc. IEEE Conf. Comput. Vis. Pattern Recog., 2012, pp. 3538–3545.
- [5]. L Neumann and J. Matas, "Scene text localization and recognition with oriented

- stroke detection,” in Proc. IEEE Int. Conf. Comput. Vis., 2013, pp. 97–104.
- [6]. A Bissacco, M. Cummins, Y. Netzer, and H. Neven, “PhotoOCR: Reading text in uncontrolled conditions,” in Proc. IEEE Int. Conf. Comput. Vis., 2013, pp. 785–792.
- [7]. A Fischer, A. Keller, V. Frinken, and H. Bunke, “HMM-based word spotting in handwritten documents using subword models,” in Proc. 20th Int. Conf. Pattern Recog., 2010, pp. 3416–3419.
- [8]. V Frinken, A. Fischer, R. Manmatha, and H. Bunke, “A novel word spotting method based on recurrent neural networks,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 34, no. 2, pp. 211–224, Feb. 2012.
- [9]. R Manmatha, C. Han, and E. M. Riseman, “Word spotting: A new approach to indexing handwriting,” in Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recog., 1996, pp. 631–637.
- [10]. T Rath, R. Manmatha, and V. Lavrenko, “A search engine for historical manuscript images,” in Proc. 27th Annu. Int. ACM SIGIR Conf. Res. Develop. Inform. Retrieval, 2004, pp. 369–376.
- [11]. T. Rath and R. Manmatha, “Word spotting for historical documents,” *Int. J. Document Anal. Recog.*, vol. 9, pp. 139–152, 2007.
- [12]. J. A. Rodriguez-Serrano and F. Perronnin, “Local gradient histogram features for word spotting in unconstrained handwritten documents,” presented at the Int. Conf. Frontiers Handwriting Recognition, Montreal, QC, Canada, 2008.
- [13]. J. A. Rodriguez-Serrano and F. Perronnin, “A model-based sequence similarity with application to handwritten wordspotting,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 34, no. 11, pp. 2108–2120, Nov. 2012.
- [14]. S. Espana-Bosquera, M. Castro-Bleda, J. Gorbe-Moya, and F. Zamora-Martinez, “Improving offline handwritten text recognition with hybrid HMM/ANN models,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 33, no. 4, pp. 767–779, Apr. 2011.
- [15]. I. Yalniz and R. Manmatha, “An efficient framework for searching text in noisy documents,” in Proc. 10th IAPR Int. Workshop Document Anal. Syst., 2012, pp. 48–52.
- [16]. K. Wang, B. Babenko, and S. Belongie, “End-to-end scene text recognition,” in Proc. IEEE Int. Conf. Comput. Vis., 2011, pp. 1457–1464.
- [17]. A. Vinciarelli and S. Bengio, “Offline cursive word recognition using continuous density hidden Markov models trained with PCA or ICA features,” in Proc. 16th Int. Conf. Pattern Recog., 2002, pp. 81–84.

Cite this article as :

Shivali Parkhedkar, Shaverri Vairagade, Vishakha Sakharkar, Bharti Khurpe, Arpita Pikalmunde, Amit Meshram, Prof. Rakesh Jambhulkar, "Handwritten English Character Recognition and translate English to Devnagari Words", *International Journal of Scientific Research in Computer Science, Engineering and Information Technology (IJSRCSEIT)*, ISSN : 2456-3307, Volume 5 Issue 2, pp. 142-151, March-April 2019. Available at doi : <https://doi.org/10.32628/CSEIT19528>
Journal URL : <http://ijsrcseit.com/CSEIT19528>