

KANNADA HANDWRITTEN WORD RECOGNITION USING R_CLUSTERING AND SUPERVISED LEARNING DISTANCE TECHNIQUES

Shakunthala BS, Kalpataru Institute of Technology
CS Pillai, ACS College of Engineering

ABSTRACT

Handwritten text line segmentation is regarded as an important test in record picture evaluation. The difficulty in today's printed text, Skewed lines, curved lines contacting and covering components, generally words or characters, among lines, and geometrical features of lines. The difficulty involved in segmenting Handwritten Documents for Indian dialects are Telugu, Tamil, and Malayalam. Manually written archives with bended and non-parallel text lines also perform segmentation and recognition testing. This work considers text line segmentation of manually authored Kannada content records using ICA.

In this preliminary step, the authors provided an improved approach for manually written content line division to the proposed technique. The proposed system includes improved text-line segmentation as well as skew estimation and the dataset is a handwritten Kannada document. The preprocessing approaches are: (i) Filtering, (ii) Grey scale conversion, and (ii) Binarization. The ESLD method is used to estimate distance between text lines and R Clustering aids in word grouping or the Connected Components. Skew estimate can also be achieved by determining the skew angle with respect to the gap. The output demonstrates that the proposed system out performs the competition.

Keywords: Preprocessing, Grey Scale Conversion, Binarization, Textline Segmentation, Skew Detection, Correction.

INTRODUCTION

The HCR (Handwritten Character Recognition) method normally identifies writers' identities. It includes the crucial phase of segmentation, in which the handwriting text is converted into lines. Handwriting text is classified into two types: offline HCR (Handwritten Character Recognition), in which writers use a pen/pencil to write on material, and online HCR (Handwritten Character Recognition). The second sort of HCR is online, in which the writers utilise a digital device, such as an electronics pen, to create. Apparently, handwriting identification is quite difficult due to the tremendously wide variety in diverse handwriting practices of different people. In recent years, a range of machine learning approaches, such as SVM, Gaussian Mixture Modeling, ANN, Fuzzy Logic, and others, have been combined for developing algorithms for both offline and online HCR. The HCR technology must include the following features:

- **Flexibility:** This means that the system must be able to handle any type of writing variation including a large number of different people's character patterns.
- **Personalization:** The OCR system should be able to recognise online handwritten forms from any type of user.
- **Effectiveness:** Online HCR devices must be space time efficient.
- **Automatic Teaching:** The OCR systems must be trained using an automatic teaching method.

The approach used to provide customisation functionality. A HCR network contains 3 critical levels: pretreatment, extraction of features, and classification. In compared to printed characters, identifying unrestrained character recognition by an HCR system is extremely difficult. This could be due to a variety of factors such as line change, size variation, difference in space between letters, as well as the existence of skew, distortion, and so on. Various academics have advocated a variety of strategies for developing an excellent handwritten recognition software. In comparison to printed papers, text line separation becomes quite difficult for handwritten manuscripts. Text characters in printed papers are normally straight and parallel, but handwritten writings include irregularly slanted and curved lines with semi inter row spacing. The procedure becomes more complex and critical due to differences in inter-line length, the presence of unequal baseline skew, overlapped and contacting text-lines. The efficiency of term segmentation is influenced by how accurate or erroneous the text-line division is, which affects the accuracy of term recognition. For text categorization, a classic technique known as global horizontal projection examination of dots has been presented. Because of irregular inter-line lengths, text-line skew variation, and overlapped and connecting elements of two successive text-lines, this method cannot be used on unconstrained handwritten textual information right away.

The handwritten identification system is divided into three stages: pre-processing, extraction of features, and identification. Skew detection and correcting, as well as letter, phrase, and line segments, occur during the initial step of pretreatment digitization. The next step is to extract unique characteristics from a previously processed text document. Character recognition is the third procedure. Among the three stages, pretreatment stands out for its improved performance in the latter steps of the handwritten identification system. It has been observed that the optimum phase of preprocessing encounters a problem in which people's handwriting is not in a single direction but is somewhat slanted. Such an angle is known as skew, and it complicates line and part of a word. To address this issue, a proper deskewing technique is needed so that the rows and letters may be split effectively and the accuracy of identification can be increased. Because Kannada script has emphasis, implementing deskewing to a handwritten Kannada text is difficult. A Kannada script contains consonants, modifications, vowels, and composite letters. Also, due of the letter's above and below modifiers, dividing the line from the picture becomes a time-consuming procedure. Due to the obvious varying character sizes and varying word gap that exists between the words, sectioning the word gets complicated as well. The experiment has taken into account the above problem of deskewing and division relevant to lines and phrases of a handwriting Kannada text.

Preprocessing is the process of transforming a source image into a binary image. It includes extraction of features, Skew identification and prevention, binarization, remove noise, and morphological processes.

- **Skew Identification and Correcting:** The Fourier transform technique aids in the correction of image skews.
- **Binarization:** Grayscale images can be translated into machine images using the Adaptive threshold approach.
- **Noise reduction:** Noise removal can be accomplished using median filtering.
- **Morphological processes:** Morphology is used to retrieve image components.

Morphological Process

- **Dilation:** It is the process through which objects develop within a picture. It also helps with probing by organising elements and expanding forms in the source images.
- **Disintegration:** In this procedure, all images at the organism's boundary line are erased. As a result, the elements in the binary image are reducing.
- **Close:** When dilatation is preceded by degradation, the operation is referred to as close.
- **Hit or Miss Transition:** It aids in the detection of shapes and patterns in pixels in both the foreground and the background.

The application of various methods to the source image results in a pre-processed picture that is fed into the OCR program's segmentation stage.

Segmentation Techniques

- **Projection Profile:** This technique aids in supplying the quantity of ON dots gathered along parallel lines. Typically, a horizontal projection characteristic is used for text line separating, which entails the following two actions:
- Pre-processing with morphological operations
- Text line extraction

For separating the line/word inside of an unconstrained handwritten Kannada manuscript, previous research supports the a deskewing algorithm. As a starting point, the proposed technique employs preprocessing, dilatation, and tagging of associated parts of the input image. Following that, words related to the text lines are categorised using a proper manufacturing. When the angle is greater than zero, operations such as spinning and skew aspect computation are used to correctly inscribe the retrieved words in the new photo without the words overlapping in the txt file. The redesigned horizontal projection pattern method for company's segments is presented in this study. Moreover, during the word segmented image, the approach recognises text lines containing consonant modifiers, as well as any interword and intracharacter interval variations.

Related Literature

Hanmandlu et al. (2007) suggested a technique for handwritten Hindi numerical identification in Devanagari script. The process relies on an exponentially membership value that has been adapted to fuzzy sets created from variables consists of normalised distances acquired to use the box approach. The exponentially membership function is changed by two system parameters that are calculated by maximising the probability according to the basis functions being equal to unity. The overall classification accuracy is calculated to be 96%. The tests, however, are carried based on a small dataset of 3500 items.

Wen et al. (2007) colleagues proposed a method for recognising handwritten Bengali numbers and its use to an automated letter sorting device for Bangladeshi Post. There are two innovative approaches provided. One strategy is picture reconstructive identification, another is directional feature extraction mixed with PCA and SVM. The findings of the two suggested technique were merged with one traditional PCA strategy to meet the performance standards of an automated letter sorting system. It has been discovered that the combined system's identification result is more trustworthy than that of a single technique. The integrated program's median recognition accuracy, margin of error, and dependability are 95.05 percent, 0.93 percent, and 99.03 percent, respectively.

Banashree & Vasanta (2007) investigate solitary Devanagari numerals identification to use a global technique based on numerical end positions data. Backward Propagation (BP) neural pathways and neural circuits composite systems are used to feed factors based on terminal data to the identification system. Cholinergic performs the categorization at a pace of 200 characters per minute, while the BP human brain performs it at a rate of 200 bits per 2.5 seconds. Experiments conducted by them demonstrate a reliability of 92-97 percent.

Rajashekararadhya & Ranjan (2009) employed a method to compute feature extraction technique depending on zone center and picture centroid to recognise various language numbers. The letter centroid is determined, and the number picture is split into N-identical zones. The length between the letter centroid and every pixels in the zone is calculated. Likewise, the zone center is determined, as is the average difference from the zone center to every pixel in the area. Classifiers based on closest neighbor and feeder forward recursive neural networks are utilised for subsequent selection and classification.

The researchers of Pal et al. (2007) developed an improved quadratic beginning at age technique for recognising off-line scribbled numerals in six prominent Indian scripts. Kannada is among the scripts. The classifier's models are extracted from the helping develop of the Cartesian coordinates of the numerals. The test image of a numeral is subdivided into blocks enabling feature calculation, and the directional characteristics are generated in each component. These units are then downstream sampled using a Gaussian filter, and the characteristics derived from the below sampled blocks are put into a modified nonlinear classifier for detection. For conclusion computation, a five-fold test set technique was applied, and they verified 98.71 percent accuracy for Kannada inscriptions obtained by doing trials from their own collected data.

Acharya et al. (2008) employed k-means to categorise the Kannada cursive integer and 10-segment string idea, source of water, vertical and lateral strokes, and completion dates as features. They claimed an overall recognition performance of 90.5 percent. Rajput & Hangarge (2007) presented an image segmentation technique for isolated Kannada handwritten identification. Regardless of image size, various digital images matching to each scribbled numeral are combined to produce patterns, which are stored as 8x8 matrices. The numerals to be recognised are matched with every pattern using a cosine similarity classifier, and indeed the best match pattern is regarded the recognised number. They acquired an average detection accuracy of 95.62 percent by executing trials on a data set sense of good.

The preceding survey focused on studies related to numerical recognition that was present in the market. Certainly, when compared to some other ways, the productivity of any of the offered systems in terms of character recognition is excellent. The tests, however, are not

being carried out on a single data point including a huge amount of data. This is owing to the lack of a common database for integers written in Indian languages as opposed to non-Indian languages (Mezghani et al., 2012).

Enhanced Supervised Learning Length and R_Clustering Technique

Overview

The study proposes an improved version of Text region separation and Skew estimating for a Handwriting Kannada manuscript. Using the ESLL (Enhanced Supervised Learning Length) and R Clustering techniques, the recommended technique finds and separate handwritten document lines, calculating the skew for the skewed paragraphs. The ESLL (Enhanced Supervised Learning Length) algorithm does gap estimating, which entails measuring the distance between any two paragraphs or letters. The related components are then grouped using the R Clustering technique. The procedures are divided into four stages: I entry of the document picture, II preprocessing, III text line separation, and IV skew calculation.

Proposed Methodology

Preprocessing

The image of the handwritten report serves as input for preprocessing. Preprocessing is accomplished through I filtering, II grey conversion, and III binarization. This procedure ensures that the classification technique is at or above par. The delivered document picture represents the picture's row and column.

Filtering: This procedure is required to remove any unwanted pixels or distortion from the paragraphs of the source handwritten picture in order to successfully separate the text lines.

Gray scale converting: The documentation image is processed in order to retrieve the text line and letters. The suggested technique is in the shape of a grayscale image as an outcome.

Binarization: Binarization is the process of converting grayscale photographs into output images by separating the foreground or message from the background details. In those other terms, if the picture's intensities value exceeds the existing threshold level, the pixels' numbers are converted to 1, otherwise to 0. Following that, sections with no text are removed from the source images by transiting it in the up, down, left, and right directions. This includes removing any white pixel regions from the picture. As results, only the picture's textual portion is received.

Text line Segmentation

The handwritten sentences are first separated into rows, which are subsequently separated into individual words by finding the CCs. Following that, there are Linked Components based on their mean width and length. The method used works by taking into account the length and width of the entire handwritten text. When there is no skewed, there must be at least one minimum value for the word's length and one maximum for the word's breadth. Following skew adjustment with approximated skew angle repeating of the same procedure, the busy areas are considered for

exact skew restoration. The expected text lines are used to compute the average breadth, length, and line length.

To accomplish this, *Imag1* is vertically divided into *n* equal portions, and the *avghlinehyt* for every line of text is computed by calculating the text lengths for the whole file picture into account. The total number of data pixels every row of every vertical division of arrays *Imag1* is used to compute the length. For accessible values, this information is entered row by row into matrix *CPY*. The zero data pixel position in array *CPY* denotes to the gap between two consecutive lines of text that aids in calculating the length for the same: Each text's *maxlinehyt* and *minlinehyt* values, as well as its placement information, are stored in distinct arrays. For unconnected diacritics, the *minwidth* and *minhyt* associated to the appearance have been established. The data is represented by the black pattern in the handwriting letter, while the background is represented by the dark one.

The problem with overlapping and linked lines of text can be handled by determining the gap between the lines of text. When a segment is received, the length is confirmed; if the length is below or equivalent to the medium height of the liner, the *txt* file is single; otherwise, if the length is more, the line drawn has two or even more texts. This holds true for all successive vertical separation, and the operation is performed for each section separately. If the section's height is smaller than the height limit and it is adjacent to another section, it is diacritic. If the breadth of the segment is less than or equivalent to the minimal after vertical separation, it shows that the two signalling are joined at a certain location. It must be lowered to the level where there is lowest amount of data loss, because erroneous text region separation can resulted in wrong word separation. This can lead to incorrect feature extraction and recognition. In cases where the linked character's location is erroneous, *maxlinehyt* or *minlinehyt* could be used to reduce segmentation problems (Rajashekararadhya & Ranjan, 2009).

Group CC Utilizing R_Clustering Technique

Elements of same line of text can be joined in a subset using the learning distance metric, but elements from other text lines can also be linked. Those between-line borders are opaque because their widths may not be greater than the widths of the within-line margins. This problem can be managed using hypervolume, in which the total of the hypervolumes of the cluster's nodes is used to compute the separation:

$$G_u = \sum_{i=1}^m [\det(C_i)]^{1/2} \quad (1.1)$$

It $[\det(C_i)]$ denote the determinant of the convolution of Cluster *I* is referred to here. It is computed using the constituents black dots of the asteroid's CCs Figure 1.

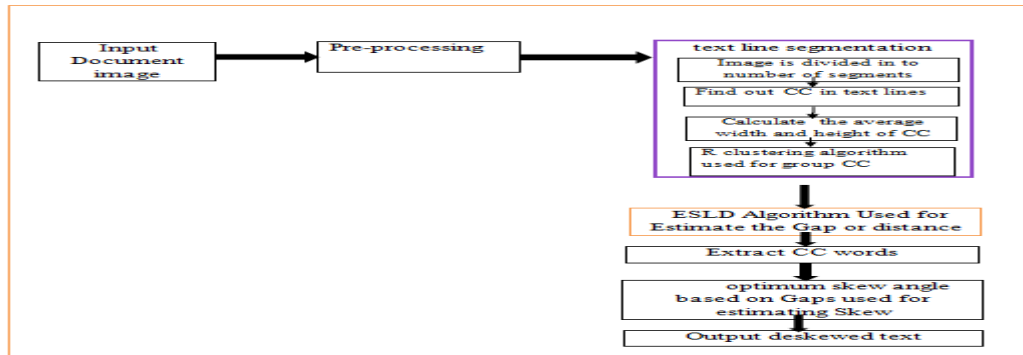


FIGURE 1
SEGMENTATION ARCHITECTURE

At first, the spanned tree's components work as a single group, then each single line is erased, dividing the groupings between two sub-clusters, referring to a discontinuous subtree. The border with the greatest reduction in measurements is picked and deleted in order to lower the total size of the file. This is known as the maximal hyper-volume decrease criteria, and it is denoted as:

$$\text{edge}_{\text{deleted}} = \arg_{\text{edge}}^{\max} \Delta G_u = \arg_{\text{edge}}^{\max} [G_u(F_k) - G_u(F_{k+1})] \quad (1.2)$$

Wherever Where $F_k = \{T_1, T_2, \dots, T_K\}$ (1.3) represent the division of k disconnected subtrees. The metric does not allow for the analysis of the number of clusters. The reason for this is that it always lowers as the number of nodes increases. Text lines could be considered to have great rectangular shapes almost all of the time. If it was curved, it had to be divided into several straighter sublines. It is expected and with an acceptable clustering (segmented texts), the score of text line flatness is the greatest. The entire straightness measurement is depicted below:

$$F_{um} = \sum_{i=1}^k \left(\frac{\lambda_{i1}}{\lambda_{i2}} \right)^2 \quad (1.4)$$

With k equal to the number of groups, and λ_{i1} and λ_{i2} ($\lambda_{i1} > \lambda_{i2}$) the Eigen values of each cluster's covariances.

The SLD Algorithm is used to Estimate Gaps or Distances

Clearly, most clustering algorithms rely heavily on a good metrics between the pairs of input nodes, therefore establishing the distance between the CCs is critical for guaranteeing the resulting spanning tree contains components with same line in a subset and separate subtrees for multiple paths. A distant metric method is aimed to carry out text line separation utilizing supervised learning and is inspired by the theory of distant metric learning. To indicate the self-designed distance metric amongst some of the CCs nodes, training data of element pairs are used

and labelled as "*inside line*" and "*beyond lines*." These are accomplished by using the truthing software TTLC to label specific training material pictures. It essentially aids in the labelling of lines of text and letters through automatic transcript matching and manual correction. Imagine a training document with a collection of clusters, such as $F=b_1, b_2 \dots b_n$, where n denotes the amount of parts. Two sets of element pairs are obtained from this, which serve as examples for data augmentation:

$$\begin{aligned} X &= \{(b_i, b_j) / b_i \text{ and } b_j \text{ belong to the same line}\}, \\ Y &= \{(b_i, b_j) / b_i \text{ and } b_j \text{ belong to different line}\} \end{aligned}$$

Because it is assured that only spatially neighbouring elements are connected in the minimum spanning tree, certain element couples can be omitted from the test set to improve metrics learning. This is accomplished by constructing the trained document's surface Voronoi chart, which represents the spatial adjacency between the elements. If the elements share a Voronoi border on their boundary, the a_i element is said to be a neighbour of the a_j element. In X and Y , the non - parallel couples in the Voronoi diagram are removed. Metric learning focuses on lowering the distance between the elements in X and amplifying or increasing the distance between the elements in Y in relation to the learnt metric. As a result, the metric learning task can be modelled as a concave programming problem.

$$\begin{aligned} \min_{B \in R^{m \times m}} \quad & \sum_{(b_i, b_j)} \|b_i - b_j\|_A^2 \quad B \geq 0, \\ \sum_{(b_i, b_j)} \|b_i - b_j\|_A^2 \quad & \geq 1 \quad (1.5) \end{aligned}$$

The length metric is called here by the matrix $B \in R^{k \times l}$ (where k denotes the characteristic gap describing the element pairs). Dimensionality

$$d(b_i, b_j) = d_B(b_i, b_j) = \|(b_i, b_j)\|_A = \sqrt{u_{ij}^T * B * u_{ij}} \quad (1.6)$$

The feature map given by describes v_{ij} the relationship between two places a_i and a_j . The concave programming issue is used to solve B .

Extraction of Connected components (CC)

Every element C_i in the collection X is an element of the series. Allow the collection of elements to be labelled as $C_i N$ within the C_i neighbourhood $N C_i$. The C_i component looks for surrounding elements that can satisfy any of the following conditions:

C_i covers at least a portion of the length of C_j ; $j \in N$, or vice versa.

The height-wise midpoints of C_i and C_j ; $j \in N$, have a vertical displacement smaller than a criterion 2 of C_i or C_j 's length.

If the element C_{ij} meets any of the criteria mentioned above, the element is allocated to the C_i line. C_i 's boundaries are initialised with the limits of C_i . In addition, the region of the

line expands based on the number of elements allocated to a certain line. Assume that a text line has all of its elements tentatively or entirely within its frame. These elements can be marked in the same way that the txt file is. Small components can be implementation of the overall organization to their relevant lines of text in this manner.

Optimal Skew Angle Skew Prediction

This process entails identifying the related components to determine the optimal skew values of CCs. Initially, the Edge Analyzer is used to preprocess and identify the borders of the captured images. The retrieved edges are then dilatable using a round transformation function to find the related components. Skew inclination estimation is accomplished by locating the centroid of each and every coupled component in the text and plotting an ellipsoid on it. The azimuth of a CC denotes the angle formed by the reference direction and the primary or major axis through which the CC circles with the least amount of inertia. Basically, it should be provided that the region's second moment is identical to the ellipsoid. The MID (K, L) of Coupled Components (CCs) is calculated as

$$MID(K, L) = \left(\frac{\sum OrigK_i}{N}, \frac{\sum OrigL_i}{N} \right) \quad (1.7)$$

Where, $(OrigK_i, OrigL_i)$ the co-ordinate quantities of pixels within associated CCs are depicted, and N indicates the total number of pixels within each Coupled Component. First order features are calculated:

$$(L_i) = (OrigL_i - L) \quad (1.8)$$

Second order variables are calculated using the preceding mathematical equations

$$\mu_{kk} = (\sum (K_i)^2 / N) \quad (1.9)$$

$$\mu_{ll} = \left(\sum \frac{(L_i)^2}{N} \right) \quad (1.10)$$

$$\mu_{kl} = \left(\frac{\sum (K_i * L_i)}{N} \right)$$

The μ_{ll} and μ_{mm} seconds and represent the l and m differences, respectively

$$\text{Skew}(\mu_{ll} > \mu_{mm}) = \frac{\mu_{mm} - \mu_{ll} + \sqrt{(\mu_{mm} - \mu_{ll})^2 + 4 * \mu_{lm}^2}}{2 * \mu_{lm}}$$

$$\text{skew}(\mu_{mm} > \mu_{ll}) = \frac{2 * \mu_{lm}}{\mu_{ll} - \mu_{mm} + \sqrt{(\mu_{mm} - \mu_{ll})^2 + 4 * \mu_{lm}^2}}$$

$$\text{Optimum skew angle } (\theta) = (180/\pi)\tan^{-1}(\text{skew})(1.11)$$

RESULTS

The suggested deskewing method is tested on handmade Kannada scanned documents with 250 line stuff and 754 letters. The averaged line segmentation provides a 97.13 percent accuracy, whereas word separation yields a 93.48 percent accuracy. It is showed that the suggested methodology achieves comparable performance in the separation of skewed line stuff and words. The SLD algorithm is used to effectively accomplish line splitting in this method. It has been discovered that the effectiveness of the term segmentation stage degrades as a result of irregular spacing within words and fractured characters. Table 1 summarises the results of line and character segmentation. The suggested method has been tested on handwritten text, including data acquired from people from various backgrounds. Approximately 4600 Kannada handmade words are used in the research, with 42 percent used for classifier training and the remainder used for testing. Approximately 2700 Kannada items have been tested, producing an accuracy of 96.15 percent Figure 2.

Table 1 OVERALL PERFORMANCE OF THE PROPOSED SEGMENTATION			
Methods	Textlines	Words	Accuracy
Input	250	754	97.13%
Proposed segmentation	241	694	93.48%

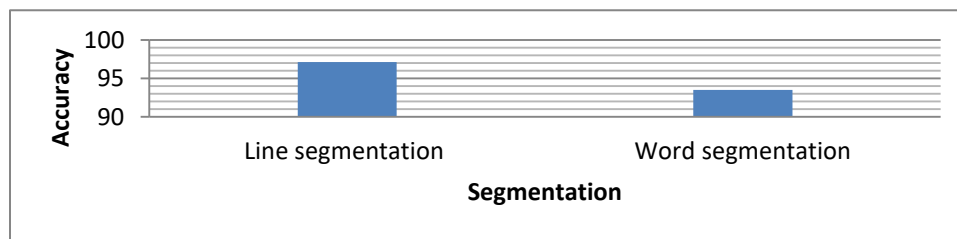


FIGURE 2
SEGMENTATION OF LINES AND WORDS

CONCLUSION AND FUTURE WORK

The first achievement is made possible by mathematical functions, which is utilised to build bridges between items in order to dealing with skewed textual and analyse enhanced horizontal line projections, diagonal elements grouping, and splitting on a separation of the input

data into vertical bands. We presented a Modified oriented gradient profile and linked component approach for text line extraction of Kannada handwritten manuscripts in this study. The approach was tested on completely unrestricted handwritten Kannada papers and achieved a separation rate of 97.5 percent on average.

The second achievement of the study work is establishing accurate and effective handwritten text line fragmentation, which is extremely difficult and in high demand due to its numerous possible uses. Kannada character recognition with application to scheduled structure management is proposed in this work. For planning, only 57 characters are evaluated. Only transcribed letters are extracted using suitable pre-planning systems. For features extraction, PCA and HOG are used.

The third goal of the suggested study approach is to estimate skew using the deskewing ESDCW technique, which results in the identification of lines and words from handmade scanned documents. A deskewing technique for line and part of a word from a printed and handwritten Karnataka document is proposed in the study effort. Following that, words related to the text lines are categorised using a proper manufacturing. Following that, the recognised words are extracted and saved in a new picture. Unwanted information is selectively deleted during word harvesting using the locality preserving approach, and crowding of words is avoided while saving in a new photo file. Furthermore, the study compares enhanced horizontal projection feature and CC (Connected Component) approaches for text line separation in Kannada handmade documents to purely uncontrolled handwritten Kannada papers. The overall average segmentation rate is typically 97.5 percent.

The fourth goal of the paper was to create unique ways to unconstrained handwritten Arabic line segmentation and deviation estimation. The current study suggested a strategy based on SLD and the R Clustering technique for line and character segmentation in the framework of a handwritten documents Kannada literature. The words from the text lines are arranged here using an intelligent technique. The words that have been detected are then retrieved and saved in a new picture. During extraction, it is assured that no words intersect and that unnecessary information is effectively removed. Other proposed and well tested approaches to speed up the process of text region segmentation and skew repair of Kannada handwritten manuscripts are given in the investigation. The claimed average segmentation accuracy is 97.13 percent.

Scope for Future Work

Future work could contrast Kannada and English documents using the provided methodologies and the same infrastructure used for unrestrained and skew prediction. There is a great demand for researchers to conduct research in the area of script recognition systems, particularly for South Indigenous languages. It can also be focused on crucial details for creating such an identification system for Tamil and Telugu. With the right architecture and dataset, a deep neural network may achieve a respectable recognition performance for handwritten documents. Future implementations could include identifying words and, later, phrases. The following level will be comprehending the sentences and providing satisfying responses. Once the network understands the words, the network can be constructed to summarise the meaning of the provided text file or to translate it into another language.

REFERENCES

- Acharya, D., Reddy, N.S., & Makkithaya, K. (2008). Multilevel classifiers in recognition of handwritten Kannada numerals. *International Journal of Computer and Information Engineering*, 2(6), 1908-1913.
- Banashree, N.P., & Vasanta, R. (2007). OCR for script identification of Hindi (Devnagari) numerals using feature sub selection by means of end-point with neuro-memetic model. *International Journal of Computer and Information Engineering*, 1(4), 883-887.
- Hanmandlu, M., Nath, A.V., Mishra, A.C., & Madasu, V.K. (2007). Fuzzy model based recognition of handwritten hindi numerals using bacterial foraging. In *6th IEEE/ACIS International Conference on Computer and Information Science (ICIS 2007)*, 309-314.
- Mezghani, A., Kanoun, S., Khemakhem, M., & El Abed, H. (2012). A database for arabic handwritten text image recognition and writer identification. In *2012 international conference on frontiers in handwriting recognition*, 399-402.
- Pal, U., Sharma, N., Wakabayashi, T., & Kimura, F. (2007). Handwritten numeral recognition of six popular Indian scripts. In *Ninth international conference on document analysis and recognition (ICDAR 2007)*, 2, 749-753.
- Rajashekararadhya, S. V., & Ranjan, P. V. (2009). Zone based feature extraction algorithm for handwritten numeral recognition of Kannada script. In *2009 IEEE International Advance Computing Conference*, 525-528.
- Rajput, G.G., & Hangarge, M. (2007). Recognition of isolated handwritten Kannada numerals based on image fusion method. In *Pattern Recognition and Machine Intelligence: Second International Conference, PReMI 2007, Kolkata, India, December 18-22, 2007. Proceedings 2* (pp. 153-160). Springer Berlin Heidelberg.
- Wen, Y., Lu, Y., & Shi, P. (2007). Handwritten Bangla numeral recognition system and its application to postal automation. *Pattern recognition*, 40(1), 99-107.

Received: 04-Mar-2023, Manuscript No. JMIDS-23-13334; **Editor assigned:** 06-Mar-2023, Pre QC No. JMIDS-23-13334 (PQ); **Reviewed:** 20-Mar-2023, QC No. JMIDS-23-13334; **Revised:** 22-Mar-2023, Manuscript No. JMIDS-23-13334(R); **Published:** 29-Mar-2023