

A COMPLETE OCR SYSTEM USING CNN - HMM HYBRID APPROACH FOR PRINTED MEITEI/MEETEI SCRIPT DOCUMENTS

YANGLEM LOIJING KHOMBA KHUMAN

Research Scholar, Department of Computer Science, Manipur University, Canchipur, Imphal West, Manipur. Corresponding Author Email: loijingyanglem@gmail.com

H. MAMATA DEVI

Assistant Professor, Department of Computer Science, Manipur University, Canchipur, Imphal West, Manipur. Email: mamata_dh@rediffmail.com

Abstract:

Scanned document image is transformed into a text document that can be edited using optical character recognition (OCR). Character recognition in Meitei/Meetei script is crucial because it has numerous uses including document digitization, bank data processing and mailing system automation. Hence various approaches have been presented for character recognition but these approaches face difficulty in recognition due to overlapping nature of Meitei/Meetei script characters. Also, these character recognition models were having sequential data handling problem due to the absence of dynamic temporal model in training process that leads to low recognition rate of Meitei/Meetei script characters. To solve these issues a novel Hybrid CNN-HMM Recognition Approach has been proposed in which HMM is used as a dynamic temporal model to train CNN thereby, eliminate the sequential data handling issue while recognizing characters of Meitei/Meetei script. Also, CNN and HMM were utilized to deal with character appearance fluctuations and provide temporal modelling in recognition of characters along with effective extraction of features from each scanned Meitei/Meetei script image with high recognition rate. The result obtained shows that the proposed hybrid model has high precision, sensitivity, recall, accuracy and F-measure when implemented in Python and are compared to other existing methodologies.

Keywords: Graphics separation, Skew correction, Horizontal projection profile, vertical projection profile, CNN, Hidden Markov Model (HMM)

1. INTRODUCTION

Sharing accurate transcriptions of the books will help the academic research community analyses the data on a big scale and gain new knowledge about the history of books and publishing in India. Up until this point, the majority of the content could only be accessed physically by going to the library. A crucial phase in the recognition procedure is page analysis, which includes text recognition, page segmentation, and region categorization. Its effectiveness has a substantial impact on a digitization system's overall success, not just in terms of OCR accuracy but also in terms of the value of the extracted information [1, 2]. Over the past few decades, there has been a lot of study done on character recognition. Character recognition is typically used to identify characters. Character recognition, however, refers to the identification of a character without the use of a person. There has been a significant amount of research in this field. The majority of these studies have focused on the recognition of Japanese, Chinese, and English characters. Indian

languages like Meitei/Meetei Script and Devanagari are given far less attention than other languages [3-5]. Despite the extensive research that has been done in this field, there is still space for advancement [6-8]. Researchers from all over the world have been studying HCR for a long time. The HCR problem is challenging to resolve due to the special characteristics of various scripts, as well as changes in writing styles among people and variations in a single person's handwriting according to their mental state when writing. A person with a calm disposition will have neat handwriting, but someone under stress is likely to write differently, which will result in different handwriting, and so on [9].

OCR has recently been a growing interest in digitizing the top to bottom dimensions of books and documents that existed before the broad use of digital technologies. Text that has been digitized may be quickly processed for a variety of purposes, and it is easier to execute operations such as searching and sorting on it [10, 11]. Offline mode offers the whole image of the character that need to be recognized as input. The magnitude of the language being studied is frequently connected with the difficulty of the recognition. Identification would be significantly more difficult if the language had a larger number of characters than if the language had a smaller number of characters. Similarly, it is necessary to evaluate how the various characters are written, as well as the contrasts between them [12, 13]. The characters always have an impact on the handwriting recognition system's performance. Since, distortions and pattern variations are the most difficult aspects of recognition, feature extraction is critical. Manually selecting features may result in inadequate information to accurately forecast the character class. However, because of the increased dimensionality, a large number of characteristics usually causes complications [14, 15].

OCR includes categorizing character images into machine-identifiable classes, which is a classic challenge in pattern recognition and artificial intelligence. Text can be found in many different types of images, including memos, whiteboards, medical records, old documents, and text typed with a stylus. A thorough OCR system must therefore enable the identification of characters in images. This demonstrates the necessity for additional study into large-scale character recognition systems for printed documents across a range of languages and scripts [16]. Numerous approaches, including dynamic programming, neural networks, expert systems, and combinations of all these techniques have been suggested for solving the challenge of numeral identification over the years [17]. Despite the recent introduction of numerous offline OCR systems, processing its documents remains difficult due to linguistic-based criticalities, a huge character set, complicated conjuncts, and the typical geometric structure of characters, zone-based form, and use of top line. In addition to these difficulties, it is consistently observed that processing unconstrained character identification is far more difficult than printed ones [18].

Today's offline recognition systems frequently use line-level methods that pair features extraction with convolutional neural networks (CNN) [19, 20], Long Short-Term Memory

(LSTM) [21], recurrent neural networks (RNN) [22] and Connections Temporal Classification (CTC) loss training [23]. These data-driven deep-learning systems learn features directly from the training data like they do in other disciplines, as opposed to conventional methods that employ hand-engineered features. Despite the fact that these techniques have significantly increased character recognition accuracy on publicly available benchmarks, expanding these systems to support additional domains or languages can be challenging due to the expense and difficulty of amassing a sizable corpus of training data made up of line images. Hence to tackle these issues in data preprocessing, feature extraction and recognition of characters, a novel solution has to be developed. Main contribution of this research is as follows:

- In Meitei/Meetei script printed document recognition, a hybrid recognition strategy has been introduced by integrating CNN-HMM to increase the recognition accuracy without changing character appearance.
- Furthermore, the issues in preprocessing such as boundary distortion, presence of graphics and skews were removed by using Image binarization with hybrid adaptive threshold, Flood-fill operation-based graphics separation and Entropy and HPP-based skew detection.
- In OCR system, the accuracy of recognition is improved by performing segmentation of overlapping characters in Meitei/Meetei script using vertical and horizontal Projection profile technique.
- The remaining content of the paper is organized as follows: section 2 describes literature survey, section 3 provides novel solution, Section 4 presents the implementation findings and its comparison, and section 5 concludes the paper.

2. LITERATURE SURVEY

Hamdan et al [24] numerous handwriting recognition approaches, including touch input from a mobile device and image files, were addressed. The recognition techniques use a variety of approaches, including artificial neural networks and statistical methods, to handle nonlinearly divisible problems. The methods for comparing and identifying characters in image documents are discussed in this work. The study report also contrasts several statistical techniques, including graphical methods, structure pattern recognition, and template matching. The dataset training for this method takes long time, and selecting the kernel function is difficult. Narayanan et al [25] investigate and implement image character recognition using CNN. The process of reading text from image data and documents and processing it for use in programmes like machine translation and pattern recognition is known as handwritten character recognition. Also, investigates the use of CNN for more accurate detection and recognition of handwritten text images. The CNN model is validated by its performance on handwritten characters. Using several layers, the model extracts features from photos. These are then used to

train the model, allowing it to recognize characters. Even though, this approach has limited successes due to the unevenly distributed images.

Inunganbi et al [26] recognize handwritten Meitei Mayek by extracting a local texture descriptor and projection histogram feature. Different iterations of the local binary pattern (LBP), such as the uniform LBP (ULBP), improved LBP (ILBP), and center symmetric LBP (CS-LBP), are used to express local texture descriptors. These attributes are submitted to machine learning techniques like k nearest neighbor (KNN), support vector machine (SVM), and random forest (RF) for character classification together with the projection histogram. The experiments of these feature descriptors with classifiers were looked at using a self-collected handwritten Meitei Mayek character dataset of 9800 samples. However, recognition of multifaceted characters set from handwritten documents was not considered in this study. Hazra et al [27] presented a distinct structure to create a CNN architecture from the ground up with the ability to combine feature extraction and classification. The mathematical justification for using non-linearity in the deep learning model is also found in this study. The two existing Bangla datasets, cMATERdb and ISI Bangla datasets, are applied using the four layers of the proposed CNN architecture. The "Mayek27" dataset, a projected dataset of Manipuri characters, also uses the same methodology. Furthermore, compare and contrast alternative batch sizes and optimization approaches across all datasets to better understand their usefulness. However, this model is unsuitable for performing special compound character recognition because it did not incorporate activation mapping to represent intermediate learning-focused regions.

Inunganbi et al [28] introduced a convolutional neural network as a handwritten Meitei Mayek recognition technique. The grey scale of a character picture is typically used for character recognition. The construction of three-channel images for each character in this research, however, took into account the related gradient direction and gradient magnitude images. This allowed us to extract more information from gradient photos for effective recognition. Grayscale pictures and gradient images (gradient direction and magnitude) were the subjects of separate research. The recognition rate of handwritten characters utilising a complicated network is not the emphasis of this work, though. Additionally, it is not thought to expand the data set's dimension to include additional diversity, such as symbols and numbers. Sanasam et al [29] the horizontal projection histogram (HPH) was used to partition text lines and words from unrestricted handwritten texts by detecting midpoints and gap trailing between lines. The HPH is used to estimate the midpoints for the first 100 to 200 columns of the entire document. The HPH block is assessed for various situations to determine the best rows to divide adjacent lines. The suggested technique is language-independent, robust to writing variance, and segments curves, touches, and skew-lines. Word segmentation is not regarded as a separate issue, but rather works in tandem with line segmentation. In order to determine the best word separator, the vertical projection Histogram (VPH) of t columns is continuously monitored

between a lines's above and below separators. On two separate datasets of various languages (Meitei Mayek and English), the algorithm is tested. To address the character recognition issue, however, it is necessary to concentrate on creating a more effective and reliable algorithm.

[24] Requires long time for training and [25] have poor recognition precision. It is challenging to recognize complicated character sets in [26].and [27] is not suitable to perform special compound character recognition. [28] Not focusing on the recognition rate of handwritten characters with a complex network and [29] need more efficient and robust algorithm to solve the problem in character recognition. Hence to solve the aforementioned issues, a novel solution has to be proposed.

3. A COMPLETE OCR SYSTEM USING CNN - HMM HYBRID APPROACH FOR PRINTED MEITEI/MEETEI SCRIPT DOCUMENTS

The characters in Meitei/Meetei script were recognized using a novel hybrid approach named as **Hybrid CNN-HMM Recognition Approach** in which the CNN and HMM were utilized to deal with character appearance variations and for temporal modelling. This integrated approach performs training and recognition directly at the word level without the need for labelling labor and effectively extracts features from each image by sliding windows with the overlapping of successive windows. Overfitting issues are also removed during training using the Viterbi method. With improved learning capacity during the training process, the Viterbi algorithm is a dynamic programming approach that determines the best posteriori probability estimate of the most likely sequence of hidden states. Also, this hybrid model significantly increases the recognition accuracy level in character recognition. Furthermore, to enhance the performance of the novel CNN and HMM approach, improvement in preprocessing and segmentation of scanned Meitei/Meetei script image is required. Hence the preprocessing is performed by Image Binarization with Hybrid Adaptive Thresholds Strategy in which two threshold values were assigned to pixels in order to change the uneven background without border distortions. Then, the graphics were separated from the image using Flood-Fill Operation based Graphics Separation in which Flood-Fill Operation effectively remove the graphics without requiring the shape of graphics. After graphics separation, skew correction of the Meitei/Meetei script document is required for successful segmentation, and it is a critical step in the preprocessing. Hence, Entropy and HPP-based skew detection has been used in which entropy associated with Horizontal Projection Profile (HPP) performs skew correction for Meitei/Meetei script and the orientation of the document is corrected by transforming the images with different skew angles and different levels of complexity. Due to the overlapping nature of characters in Meitei/Meetei script, character segmentation is difficult. Hence, a Projection Profile technique-based Segmentation is used in which lines, words and characters were segmented separately based on vertical and horizontal projection profiles that segments overlapped neighboring characters accurately. Hence,

the proposed approach solves the issues in preprocessing, segmentation, feature extraction and recognition of characters with high recognition accuracy.

Figure 1: Block diagram of the proposed OCR system

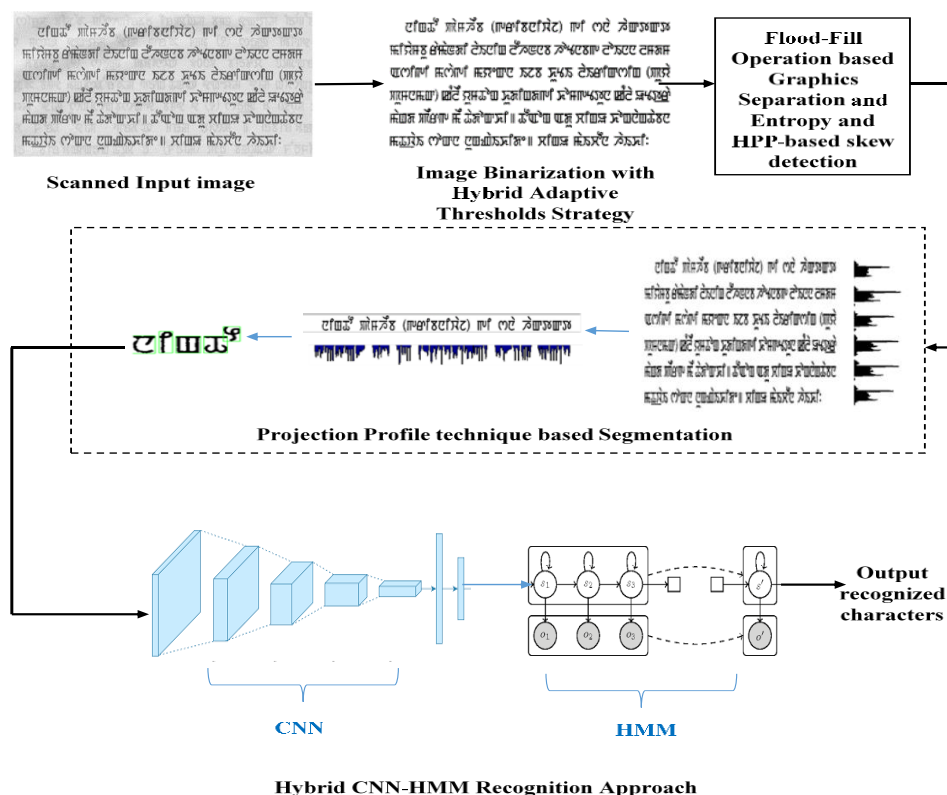


Figure 1 shows the hybrid character recognition approach in which the scanned input Meitei/Meetei script document image is preprocessed using image binarization with hybrid adaptive strategy to remove boundary distortions and blurriness. Then, Flood-fill operation based graphics separation as well as Entropy and HPP-based skew detection techniques have been included in the preprocessing step to detect and correct skews along with the separation of graphics in the input image. This preprocessed image is segmented using Projection profile technique based segmentation in which lines were segmented based on horizontal projection and words as well as characters were segmented based on vertical projections. Then, these segmented characters were given as an input to hybrid CNN-HMM recognition approach that recognizes the characters more accurately.

3.1 Hybrid CNN-HMM Recognition Approach

Hybrid CNN-HMM Recognition Approach has been proposed to recognize the characters in Meitei/Meetei script in which to improve the recognition accuracy without computational complexity due to the uneven background, noise and graphics along with improper skew. Hence these complexities are removed in the proposed hybrid approach by enhanced preprocessing processes. Then, the touched characters were effectively segmented to provide the characters as an input to the hybrid model.

3.1.1 Preprocessing

In preprocessing, the blur, distorted borders, graphics and improper skew from image were removed by dividing the preprocessing in three phases namely binarization, graphics removal and skew detection. These three phases of preprocessing were explained in the following subsections.

3.1.1.1 Image Binarization

Hybrid adaptive thresholding algorithm was used to binarize the scanned grayscale document images of Meitei/Meetei script. This system is based on two adaptive threshold values, t_L and t_G , which are derived using the local and global means of the pixels in an input picture. t_L is based on the local mean M_L of pixels within a 3×3 block size, whereas t_G is based on the global mean M_G of the entire picture's pixel. Both thresholds may or may not be employed while binarizing an input image, depending on the switched-mode: local, global, or hybrid. Each pixel has two different threshold values based on global and local variance. One of the threshold values should be picked for the local contrast region. The local standard deviation and variances of pixels inside a 3×3 block are used to alter local contrast. The steps involved in Image Binarization with Hybrid Adaptive Thresholds Strategy are as follows.

- 1) Take the input image $i_{M \times N} \in [0,1]$ which is in normalized form.
- 2) The global mean is calculated using equation (1) which is given as

$$M_G = \frac{1}{M \times N} \sum_{j=1}^M \sum_{k=1}^N i(c, d) \quad (1)$$

- 3) Calculate the local mean of neighboring pixels of $i(c, d)$ within a 3×3 block, with $i(c, d)$ at the center is given in equation (2) as

$$M_L = \frac{1}{9} \sum_{p=c-1}^{c+1} \sum_{q=d-1}^{d+1} i(p, q) \quad (2)$$

- 4) The mean derivatives of $i(c, d)$ is computed in equation (3) and (4) as follows

$$\partial_G(c, d) = L(c, d) - M_G(c, d) \quad (3)$$

$$\partial_L(c, d) = L(c, d) - M_L(c, d) \quad (4)$$

- 5) Calculate the pixels' local standard deviation within the 3×3 block and which is given in equation (5)

$$\delta(c, d) = \frac{1}{G} \sqrt{\sum_{P=c-1}^{c+1} \sum_{Q=d-1}^{d+1} \{i(P, Q) - M_L(c, d)\}^2} \quad (5)$$

- 6) The two threshold values were computed using equation (6), (7) and is given by,

$$t_G = M_G(c, d)\{1 + R(\partial_G(c, d) - 1)\} \quad (6)$$

$$t_L = M_L(c, d)\{1 + R(\partial_L(c, d) - 1)\} \quad (7)$$

Where, $R \in [0, 0.5]$ which is a bias to control the removal of background in images.

- 7) The suitable threshold value is selected using equation (8) as follows:

$$t(c, d) = \begin{cases} t_G(c, d) & \text{if } |\partial_G(c, d)| < R_1 \text{ and } \delta(c, d) < R_1 \\ t_L(c, d) & \text{otherwise} \end{cases} \quad (8)$$

Where $R_1 \in [0, 0.5]$ is a constant that controls the selection of the threshold?

- 8) Transform the pixel $i(c, d)$ into a binarized image using the threshold settings chosen and is expressed in equation (9) as:

$$B(c, d) = \begin{cases} 1; & \text{if } t(c, d) < i(c, d) \\ 0; & \text{otherwise} \end{cases} \quad (9)$$

- 9) Repeat the steps 3 to 8 until the last pixel in the image has reached.

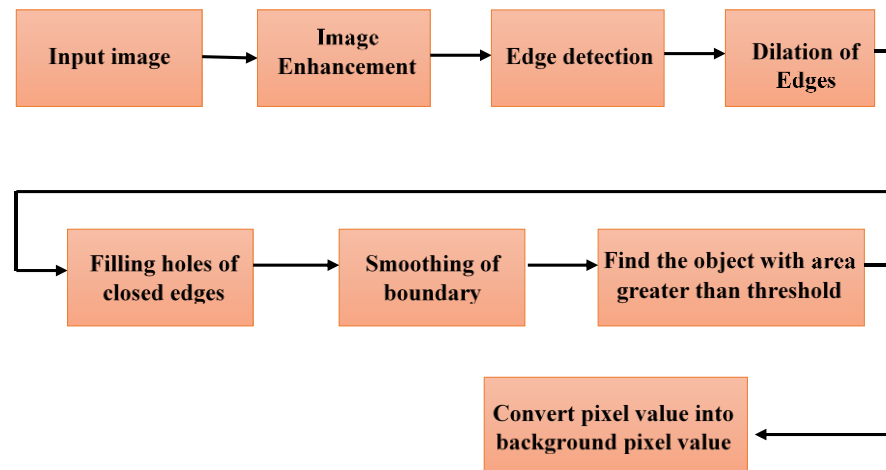
This technique is a hybrid adaptive binarization approach that adapt to the local environment to obtain a correct pixel threshold value. It includes having two dynamic threshold values for each pixel. The kind of the pixel's neighbors determines which one is chosen. Thus, this technique is utilised as a local, global, or hybrid thresholding approach. As a result, this process typically yields good results, even on extremely degraded document and thereby eliminates boundary distortion in the input image. After performing the binarization in preprocessing, it is required to separate and correct the graphics and skew from the image respectively to make the preprocessing step effective for character recognition.

3.1.1.2 Graphics Separation

For upcoming operations in information processing systems like OCR, such as segmentation, feature extraction, classification, and so on, only document images free of graphic artifacts are needed. A document image is needed to eliminate graphic objects for further processing. To make following processes like segmentation, feature extraction, and classification easier, the goal of graphical removal from a Meitei/Meetei script image

is to obtain the text region only. The initial stage in this technique is edge detection because all objects represented in the document image must have edge information for their borders. Then, all of the gaps in the objects' closed bounds are filled using the flood-fill technique. Consider the elements with a larger area to be graphics and those with a smaller space to be text. The text remains in the foreground while the pixels from the graphics area are changed to the background [30]. The process involved in flood-fill operation based graphics separation is given in figure 2.

Figure 2: Process flow of Flood-Fill Operation based Graphics Separation approach



The image enhancement is accomplished during the image binarization preprocessing step, and the edges were found using the Sobel edge detection method, which creates a gradient using a 3×3 neighbourhood of distinct dissimilar rows and columns, where the middle pixel component is weighted by 2 in each row or column to generate smoothing. The next step is dilation which involves "growing" or "thinning" the objects in a document image. The precise degree of thickening is controlled by a form provided as a structure component. A common practice in image processing called dilation allows the structural component, which is often smaller than the image, to be the second operand and the image to be the first. On the objects in the supplied Meitei/Meetei script image, a flood-fill operation has been carried out starting from the places provided in position. The marker segment and the mask picture serve to distinguish each of the applications for this morphological restoration.

Assume that binary image as b and that the image marker a , is set to 0 everywhere except for the object border, where it is set to 1 – i and is given in equation (10) as:

$$a(x, y) = \begin{cases} 1 - i(x, y) & \text{if } (x, y) \text{ is on border of } b \\ 0 & \text{otherwise} \end{cases} \quad (10)$$

Then, boundary smoothing is achieved which is a process of reconstructing an image's boundary by removing artefacts that hit it. The key to achieve the intended impact is choosing the right marker. Assume that the marker f is identified in equation (11) as:

$$f(x, y) = \begin{cases} i(x, y) & \text{if } (x, y) \text{ is on border of } i \\ 0 & \text{otherwise} \end{cases} \quad (11)$$

From the original image i , an image containing only the border object denoted as h is formed which is given in equation (11) as:

$$h = R_1(f) \quad (12)$$

Finally, the area of object in the image which is greater than the threshold is found out in order to convert the pixel value of artifacts from 0 to reference value 1. Thereby, the graphics from the image is separated separately without the need of graphic shapes. However, it is essential to correct the skew in the image to perform better character recognition.

3.1.1.3 Skew detection

Meitei/Meetei script image without graphics is given to entropy and HPP based skew detection technique that extracts the black text pixel for observation. Entropy and HPP calculations were made in order to determine the skew angle. The total of the pixel values in each vector element of a horizontal projection profile of an image is stored [31]. The HPP of the graphics removed image is calculated using equation (13) which is given as:

$$HPP(j) = \sum_{k=1}^C i(j, k) \quad (13)$$

Where, C is the number of columns and $i(j, k)$ is the image.

The entropy is a measure of randomness in statistics and is expressed in equation (14) as follows:

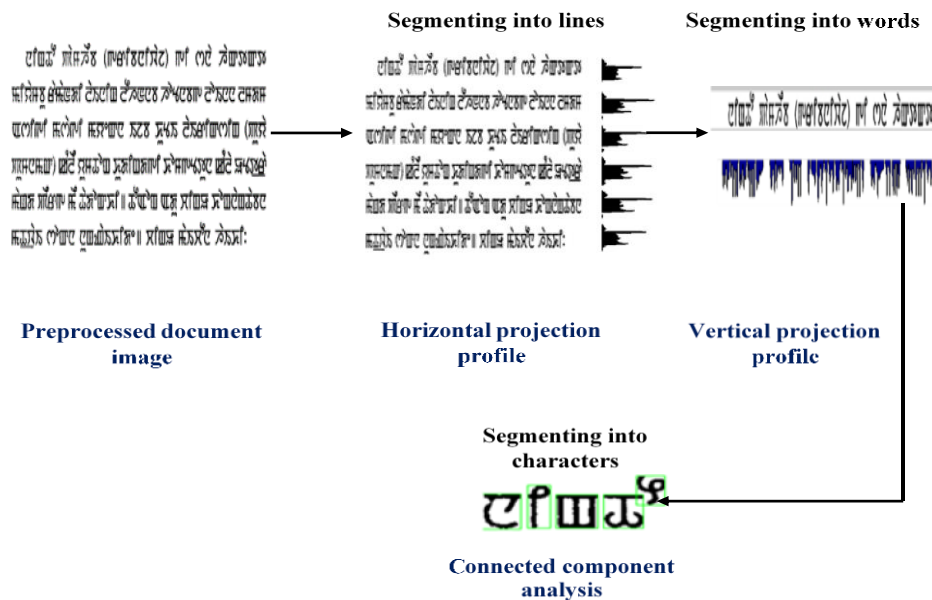
$$E(j) = \sum_j -HPP(j) * \log(HPP(j)) \quad (14)$$

The entropy associated with HPP of repeating objects is less than that of random objects like visual images or skewed text blocks. i.e., as the unpredictability rises so does the entropy. Based on the identified skew angle which is determined by entropy and HPP, the skew of the input image has been corrected. The crucial steps in preprocessing have been completed with binarization, graphics separation and skew correction. Although, the flawless segmentation of individual characters will determine the accuracy of the OCR system. Hence, the next subsection explains about the segmentation process in character recognition.

3.1.2 Segmentation

Projection profile technique based segmentation has been utilized in which preprocessed scanned Meitei/Meetei script image is segmented into lines using horizontal projection profile, the lines are segmented into words characters using vertical projection profile and words are segmented into isolated character using connected component analysis which is shown in figure 3.

Figure 3: Horizontal and Vertical Projection Profile techniques for line and character Segmentation



The steps involved in converting the image into line segments using horizontal projection profiles are as follows:

- Construct the image's horizontal projection.
- Subtract the total number of columns from each sum value in the column matrix.
- In the column matrix, find the non-zero values and modify them to 1s.
- Separate the 1s and -1s by 0s. These 1s and -1s in the image show the beginning and end of lines.
- Create a matrix using the segmented lines.
- Put the segmented line image in a folder for later use.

The steps involved in converting the line segments obtained from horizontal projection profile into words using vertical projection profile are as follows [32]:

- Construct the image's Vertical Projection.
- Subtract the total number of rows from each sum value in the row matrix.
- In the row matrix, find the non-zero values and modify them to 1s.
- Separate the 1s and -1s by 0s. These 1s and -1s mark the beginning and end of lines, respectively.
- Create a matrix using the segmented words.
- Place the segmented word images in a folder to be processed later.

Then, the characters are segmented separately using connected component analysis, the steps involved are as follows:

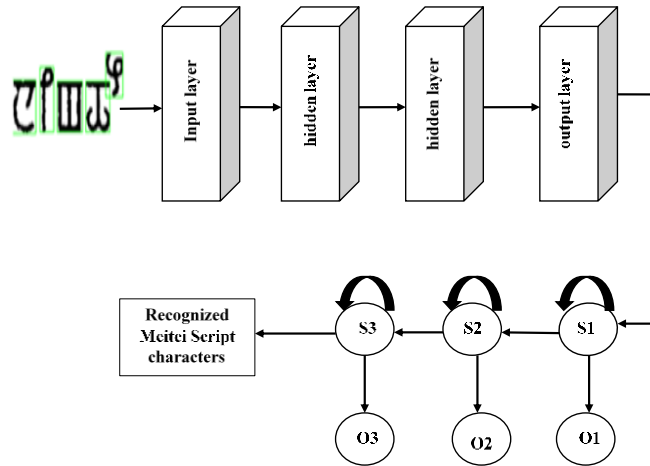
- For each unlabeled pixel, extract the related set.
- An image is divided into segments by a set of connected components.
- Draw a Bounding Box around the connecting component.
- Extract the connected component's individual character image.

Hence, the scanned document image of Meitei/Meetei script which is preprocessed has been segmented into lines, words and characters using projection profile technique and connected component analysis based segmentation and this segmentation process is suitable for overlapped characters. Then, the recognition of characters from the segmented image is important in OCR system which is explained in the next subsection.

3.1.3 Proposed Hybrid Meitei script character Recognition Approach

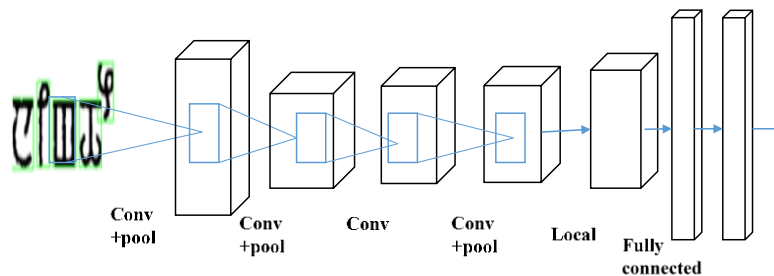
In a hybrid CNN-HMM recognition strategy, deep CNN models the shape of the pixels while HMM models the dynamics of the pixel sequence. The emission probability in characters is approximated using CNN to extract features. The Hybrid CNN-HMM character recognition model's whole architecture is shown in Figure 4.

Figure 4: Hybrid CNN-HMM character recognition model



By sliding windows that overlap one another, CNN extracts frames from each image. Then, several feature extraction techniques are used to each frame. For CNN features, the activation outputs of the first fully connected layer that is the layer preceding the output layer are retrieved. As a result, a 1024-dimensional CNN feature represents each frame. Use PCA as a feature extraction technique on the training data frames. As a result, the dimension-reduced versions of the original features are represented in the HMM observations. Four convolutional layers, one locally connected layer, and two fully connected layers make up the CNN architecture in hybrid CNN-HMM. Local contrast normalization is applied to all convolutional layers, and three convolutional layers have successive pooling layers. 11-way Softmax layer is placed in the top layer and in this CNN architecture, all connections from one layer to the next layer were connected in feedforward manner.

Figure 5: Architecture of CNN in character recognition of Meitei/Meetei script



The architecture of CNN in hybrid CNN-HMM character recognition model has been shown in figure 5. Local response normalization is included in each convolutional layer across maps to extract feature maps from characters in Meitei/Meetei script. 2×2 is the maximum pooling window size assigned in the pooling layer. At each layer, the stride alternates between 2 and 1. Rectifier units are present in all convolution and locally connected layers. Forced alignments of images were performed after training the hybrid CNN-HMM model. Each frame is assigned to an HMM state by this forced alignment. Combining the three states of an HMM into a single state makes CNN resistant to fluctuations within a category while maintaining discrimination between categories. Observation frames corresponding to these states then fall into the same category of this state. Up until the validation error stops dropping, the network is trained using stochastic gradient descent (SGD) with a cross entropy objective.

State posterior probability $P(S_T|O_T)$ is estimated by the softmax layer, but probability of observation in HMM is given in equation (15) as,

$$P(O_T|S_T) = \frac{P(S_T|O_T)P(O_T)}{P(S_T)} \quad (15)$$

Where, $P(S_T)$ is the prior probability of each state derived from the training set, and $P(O_T)$ is unaffected by the genuine digit label and thus be discarded. In a hybrid model, the output probability is the posterior probability divided by the prior probability, $\frac{P(S_T|O_T)}{P(S_T)}$ which is known as scaled likelihood.

Huge number of labelled frames was obtained after the bootstrap training phase. These data are separated to create CNN's training and validation datasets. The transition probabilities between initialized states are directly borrowed. The incorporated Viterbi method is then used to train the Hybrid CNN-HMM recognition model. Algorithm 1 summarizes the training approach.

Algorithm 1: Embedded training Algorithm in Hybrid CNN-HMM recognition model

Initialize $T \leftarrow 0$; $\Delta ap \leftarrow inf$

Train model and assign each output of HMM a label-ID.

Evaluate average precision ap_t with the recognition of each character image using Viterbi-decoding algorithm

While $\Delta ap > 0$ do

Assign each frame a label-ID as category using force alignment algorithm.

Train CNN using labelled frames

Take the prior probability $P(C_i)$ of category C_i where $i = 1, 2, \dots, n_c$

Feed each frame to CNN to get its posterior probability $P(C_i|X_j)$

Calculate scaled likelihood $P_{scaled}(X_j|C_i)$

Perform embedded training to re-estimate the transition probability of HMM and denote Hybrid CNN and HMM

Recognize each character and evaluate average precision using Viterbi- decoding algorithm

ap_{T+1}

$\Delta ap \leftarrow ap_{T+1} - ap_T$

$T = T + 1$

End

Output recognized character has been obtained from the output layer of hybrid CNN and HMM model.

Until the average precision stops increasing, this technique alternately updates HMM and CNN. It is demonstrated that the method converges to a local optimum by iteratively improving each frame's best assignment. As a result, hand-labeled frame data are not required when using supervised models like CNN. The hybrid CNN-HMM model extracts the features through the convolutional layers and update the feature map in the hidden markov layers hence, based on these training and extracted features, each character from the scanned document image of Meitei/Meetei script has been recognized accurately.

Overall, a complete OCR system using Hybrid approach using CNN - HMM for printed documents in Meitei/Meetei script has been presented to recognize the characters from the Meitei/Meetei script by removing the noise, graphics and skew from the input scanned document image as the preprocessing step. Then, segmentation was performed to separate lines, words and characters from the image and features were extracted from these images to recognize the characters using hybrid CNN-HMM recognition model. The next section explains the result obtained from a complete OCR system using Hybrid approach using CNN - HMM for printed documents in Meitei/Meetei script and discusses it in detail.

4. RESULT AND DISCUSSION

This section includes a thorough discussion of the implementation results, as well as the performance of the proposed system and a comparison section to ensure that the proposed system is applicable for OCR system.

4.1 Dataset Description

The Manipuri language is the official language of Manipur in India. The Tibeto-Burman family of languages includes this language. The dataset [33] includes binarized images, text files, and XML files for each raw image, as well as scanned 824 pages of printed papers. There are additional 51,460 isolated character samples, with 27 consonants, 7 half-consonants, 8 vowels, and 10 number characters which are shown in table 1-4. This dataset could be used in a variety of natural language processing research areas, not just for OCR.

Table 1: Consonant in Meitei/Meetei Script

ᱠ	ᱡ	ᱣ	ᱤ
ᱥ	ᱦ	ᱧ	ᱨ
ᱩ	ᱪ	ᱫ	ᱬ
ᱭ	ᱮ	ᱯ	ᱰ
ᱱ	ᱲ	ᱳ	ᱴ
ᱵ	ᱶ	ᱷ	ᱸ
ᱹ	ᱺ	ᱻ	

Table 2: Half Consonant in Meitei/Meetei Script

ᱠᱡ	ᱢᱣ	ᱤᱦ	ᱧᱨ
ᱩᱪ	ᱫᱬ	ᱭᱮ	ᱯᱰ

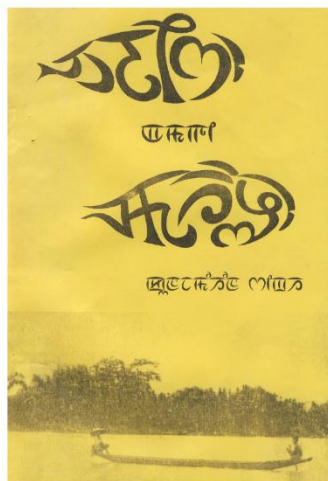
Table 3: Vowels in Meitei/Meetei Script

ᱠᱡ	ᱢᱣ	ᱤᱦ	ᱧᱨ
ᱩᱪ	ᱫᱬ	ᱭᱮ	ᱯᱰ

Table 4: Numerals in Meitei/Meetei Script

ᱠ	ᱡ	ᱣ	ᱤ	ᱥ
ᱦ	ᱧ	ᱨ	ᱩ	ᱪ

Figure 6: Sample input image



The sample input image from Manipuri language dataset has been shown in figure 6. This sample input image has the special symbols, isolated characters and numbers. This dataset is collected by scanning the book page since the background of the image seems to be uneven.

4.2 Simulated output of the proposed model

The scanned document image has been preprocessed, segmented and recognized using the proposed techniques and the output obtained from each techniques were explained in this section.

Figure 7: Binarization of input image



The uneven background and page distortion in input image has been removed by Binarization process which is shown in figure 7. Hybrid adaptive thresholding algorithm was used to binarize the scanned grayscale document image with adoption of two threshold values based on local and global mean of each pixel. Figure 7 shows the binarized image without noise and page border distortions.

Figure 8: Graphics separation in input image



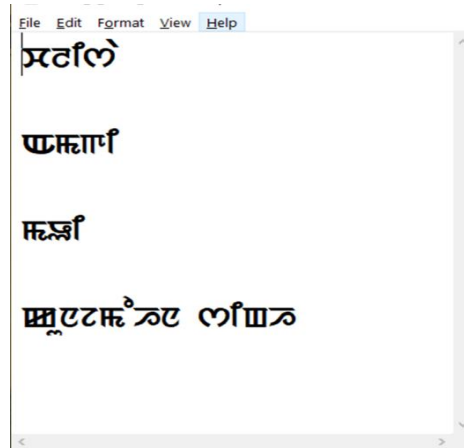
The graphics separation from the binarized input image has been depicted in figure 8. Flood-Fill Operation based Graphics Separation was used to separate the graphics which utilized Sobel edge detection to detect object edges and fill the holes in closed edges. Then, detect the area with higher threshold to convert that pixel regions into background.

Figure 9: Skew detection and segmentation of preprocessed image



Figure 9 shows the skew corrected image with segmentation results in which the skew is corrected by determining skew angle based on entropy and horizontal projection plane. Then, segmentation was performed on the skew corrected image. The lines were first segmented from the image using horizontal projection plane and the words as well as characters were then segmented using vertical projection plane.

Figure 10: Simulated recognized text output

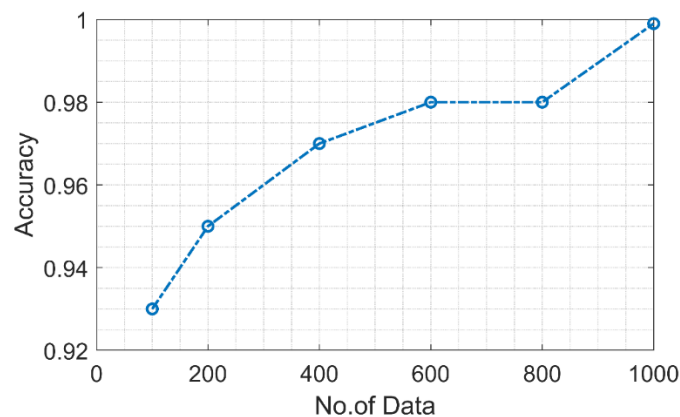


The simulated recognized text output of the characters in Meitei/Meetei script has been shown in figure 10. The features from the segmented images were extracted using hybrid CNN-HMM model with 4 convolutional layers, 1 locally connected layer, 2 fully connected layer and hidden markov chain. Through embedded training and extracted features, the characters were recognized and converted into text format with high recognition accuracy.

4.3 Performance metrics of the proposed system

The performance of the proposed approach and the achieved outcome were explained in detail in this section.

Figure 11: Accuracy of proposed model



The overall accuracy of the proposed system is shown in figure 11 and the accuracy was determined by varying the number of data from 100 to 1000. The proposed model attains the maximum accuracy value of 99.9% when the number of data is high. Also, the proposed model attains a minimum accuracy value of 93% with reduced number of data. The accuracy of the proposed system is improved by using hybrid CNN-HMM recognition model which recognize the characters without overfitting and manual data frames.

Figure 12: Recall of proposed model

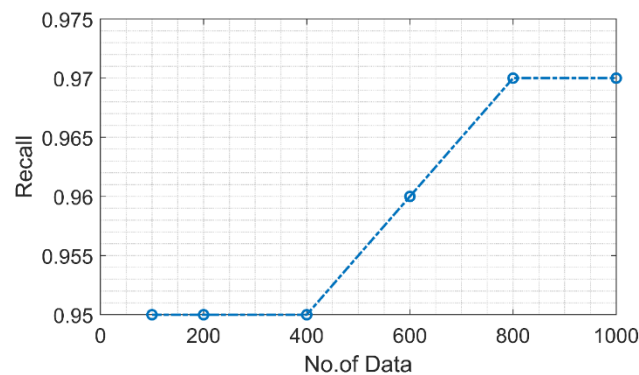
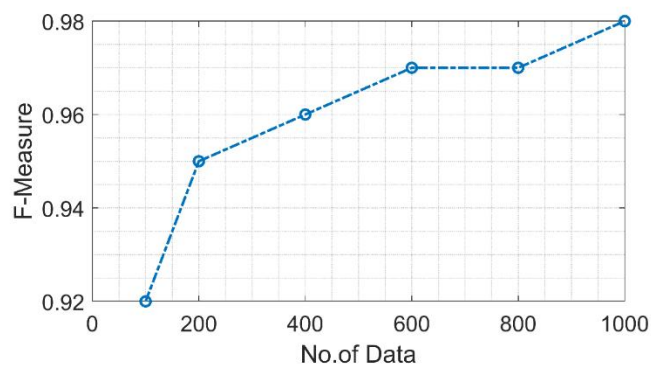


Figure 12 depicts the overall recall of the proposed model and the recall was calculated for data values from 100 to 1000. The proposed model has the maximum recall of 97% and the minimum recall value of 95%. The recall of the proposed model increases with the increase in number of data values. The recall of the proposed model is improved by hybrid CNN-HMM recognition model which efficiently extract features via the four convolutional layers in CNN.

Figure 13: F-measure of proposed model



The overall F-measure of the proposed system is shown in figure 13 and the F-measure was determined by varying the number of data from 100 to 1000. The proposed model attains the maximum F-measure value of 98% when the number of data is high. Also, the proposed model attains a minimum F-measure value of 92% with reduced number of

data. F-measure of the proposed system is improved by enhancing the recognition accuracy using hybrid CNN-HMM model.

Figure 14: Precision of proposed model

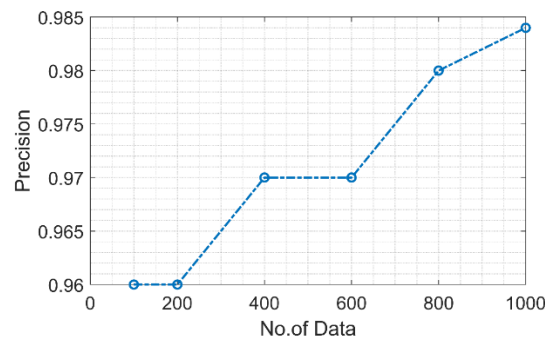
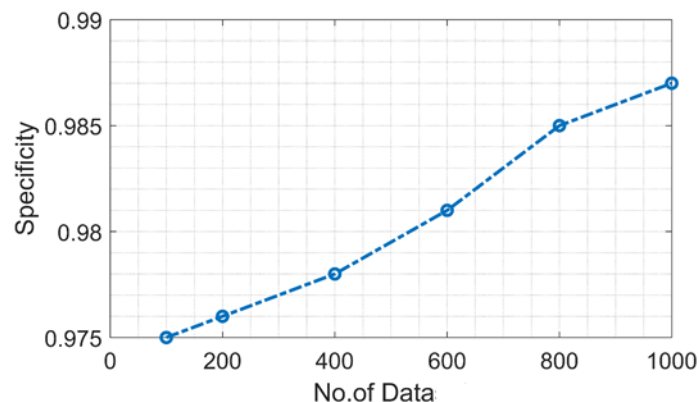


Figure 14 depicts the overall precision of the proposed model and the precision was calculated for data values from 100 to 1000. The proposed model has the maximum precision of 98.4% and the minimum precision value of 96%. The precision of the proposed model increases with the increase in number of data values. The precision of the proposed model is improved by hybrid CNN-HMM recognition model which utilizes embedded training strategy that alternatively updates HMM and CNN until the precision value stops improving.

Figure 15: Specificity of proposed model



The overall specificity of the proposed system is shown in figure 15 and the specificity was determined by varying the number of data from 100 to 1000. The proposed model attains the maximum specificity value of 98.65% when the number of data is high. Also, the proposed model attains a minimum specificity value of 97.5% with reduced number

of data. Specificity of the proposed system is improved using hybrid CNN-HMM model in which the training of the hybrid model is boosted by Viterbi decoding algorithm.

Figure 16: Sensitivity of proposed model

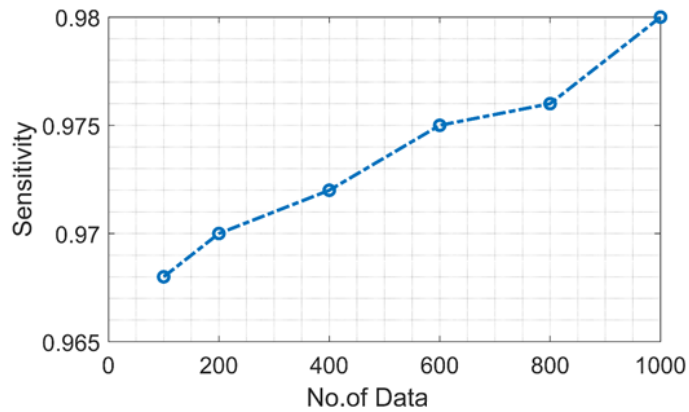
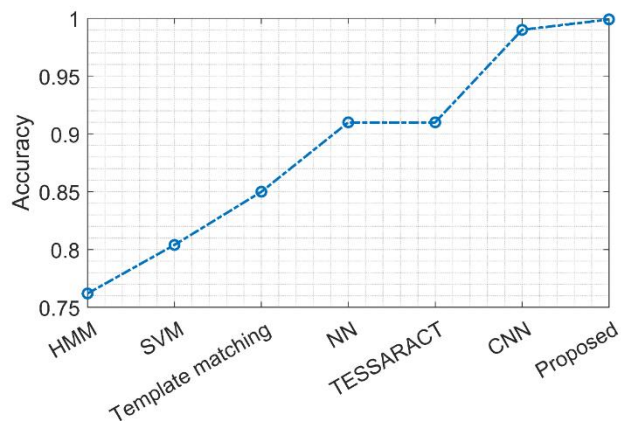


Figure 16 depicts the overall sensitivity of the proposed model and the sensitivity was calculated for data values from 100 to 1000. The proposed model has the maximum sensitivity of 98 % and the minimum sensitivity value of 96.8%. The sensitivity of the proposed model increases with the increase in number of data values. The sensitivity of the proposed model is improved by hybrid CNN-HMM recognition model which utilizes the strength of both CNN and HMM in recognizing characters without character appearance variation.

4.4 Comparison results of the proposed method

This section highlights the proposed model performance by comparing it to the outcomes of existing approaches and showing their results based on various metrics.

Figure 17: Comparison of Accuracy



The comparison of accuracy of proposed system with existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN is shown in figure 17. The accuracy of proposed model attains the maximum value of 99.9% whereas the accuracy of existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN are 76.2%, 80.4%, 85%, 91%, 91% and 99% respectively. Hence, the accuracy of proposed model is high whereas the accuracy of HMM is low.

Figure 18: Recall comparison

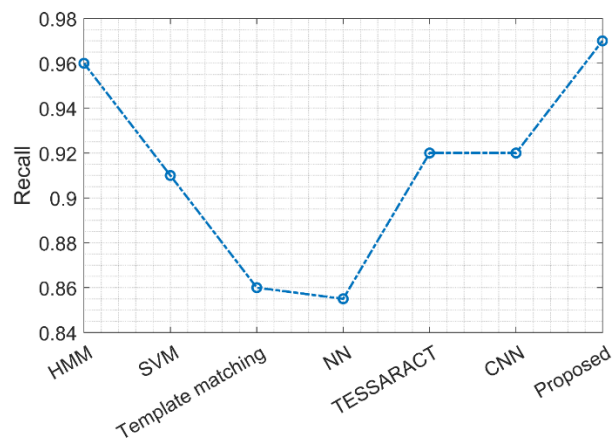
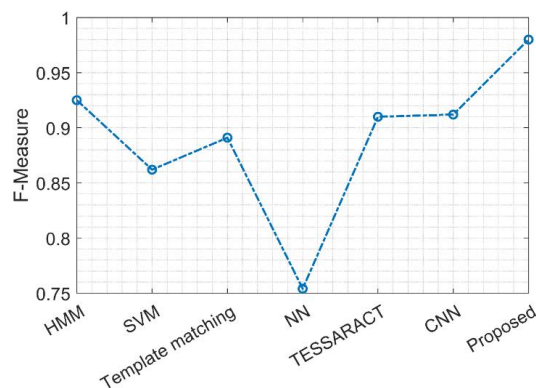


Figure 18 shows a comparison of the recall of proposed model with existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN. The recall of proposed system attains a maximum value of 97% whereas the recall of HMM is 96%, SVM is 91%, Template matching is 86%, NN is 85.5 %, TESSARACT is 92% and CNN is 92%. Hence the proposed model has the highest recall, whereas NN has the lowest recall.

Figure 19: Comparison of F-measure



The comparison of F-measure of proposed system with existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN is shown in figure 19. F-measure of proposed model attains the maximum value of 98% whereas the F-measure of existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN are 92.5%, 86.2%, 89.1%, 75.4%, 91% and 91.2% respectively. Hence, the F-measure of proposed model is high whereas the F-measure of NN is low.

Figure 20: Comparison of Precision

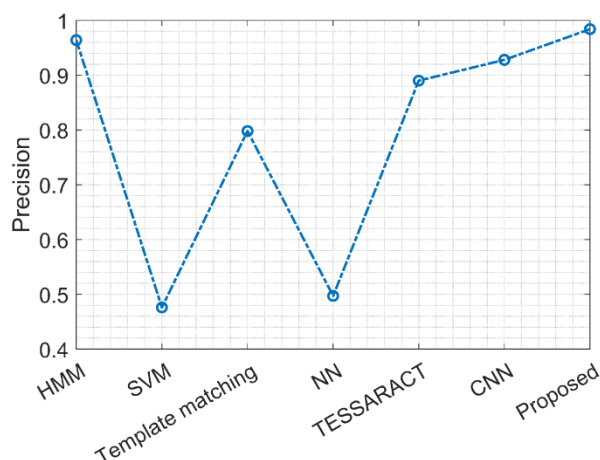
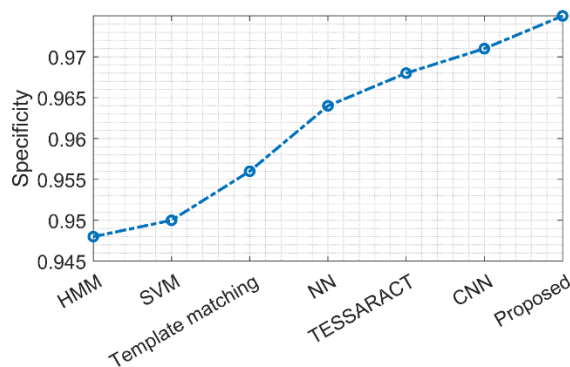


Figure 20 shows a comparison of the precision of proposed model with existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN. The precision of proposed system attains a maximum value of 98.4% whereas the precision of HMM is 96.4 %, SVM is 47.6 %, Template matching is 79.8 %, NN is 49.7 %, TESSARACT is 89% and CNN is 92.8%. Hence the proposed model has the highest precision, whereas SVM has the lowest precision.

Figure 21: Comparison of specificity



The comparison of specificity of proposed system with existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN is shown in figure 21. The specificity of proposed model attains the maximum value of 97.5% whereas the specificity of existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN are 94.8%, 94.9%, 95.6, 96.4%, 96.8% and 97% respectively. Hence, the specificity of proposed model is high whereas the specificity of HMM is low.

Figure 22: Comparison of sensitivity

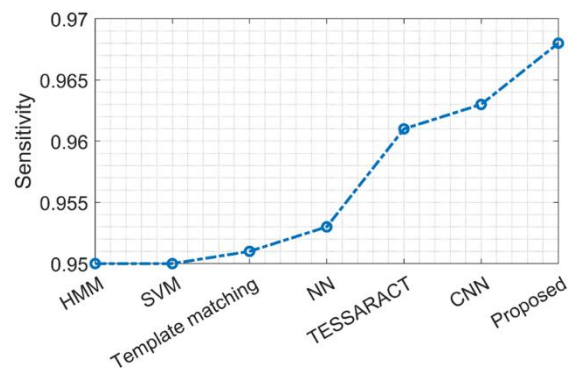


Figure 20 shows a comparison of the sensitivity of proposed model with existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN. The sensitivity of proposed system attains a maximum value of 96.8% whereas the sensitivity of HMM is 95 %, SVM is 95 %, Template matching is 95.1 %, NN is 95.4 %, TESSARACT is 96.1% and CNN is 96.8%. Hence the proposed model has the highest sensitivity, whereas SVM and HMM has the lowest sensitivity. The recognition counts of words, characters and lines from the scanned input images have been depicted in table 1.

Table 5: Recognition accuracy of words, characters and lines in scanned input images

Sample input images from scanned document	Word	character	line
Image-1	0.97	0.98	0.97
Image-2	0.95	0.95	0.95
Image-3	0.97	0.95	0.96
Image-4	0.98	0.98	0.98
Image-5	0.98	0.95	0.87

Overall, the complete OCR system using Hybrid approach using CNN - HMM for printed documents in Meitei/Meetei script outperforms existing techniques such as HMM, SVM, Template matching, NN, TESSARACT and CNN with high accuracy of 99.9%, high

precision of 98.4%, high recall of 97% and high F-measure of 98% using hybrid CNN-HMM recognition approach.

5. CONCLUSION

A complete OCR system using Hybrid approach using CNN - HMM for printed documents in Meitei/Meetei script has been proposed in this research to solve the issues in recognition of overlapped Meitei/Meetei script characters and also enhancing the preprocessing and segmentation processes to enhance recognition performance. The recognition accuracy level is improved using Hybrid CNN-HMM Recognition Approach due to the incorporation of dynamic sequence model and Viterbi training algorithm and then, the performance of this hybrid approach is further improved by image binarization, graphics separation, skew detection and projection profile segmentation. Hence, the Meitei/Meetei script characters to the recognition phase were obtained without any noise, graphics, and improper skew and with accurate segmentation of overlapped characters. The result obtained from complete OCR system using Hybrid approach using CNN - HMM for printed documents in Meitei/Meetei script outperforms existing techniques with high accuracy of 99.9%, high precision of 98.4%, and high recall of 97% and high F-measure of 98%.

REFERENCES

- 1) Binmakhashen, M. Galal, and Sabri A. Mahmoud, "Document layout analysis: a comprehensive survey," ACM Computing Surveys (CSUR), vol. 52, no. 6, pp. 1-36, 2019.
- 2) Clausner, Christian, Apostolos Antonacopoulos, Tom Derrick, and Stefan Pletschacher. "ICDAR2019 Competition on Recognition of Early Indian Printed Documents-REID2019." In 2019 International Conference on Document Analysis and Recognition (ICDAR), IEEE, pp. 1527-1532, 2019.
- 3) Lincy, R. Babitha, and R. Gayathri. "Optimally configured convolutional neural network for Tamil Handwritten Character Recognition by improved lion optimization model." Multimedia Tools and Applications, vol. 80, no. 4, pp. 5917-5943, 2021.
- 4) Memon, Jamshed, Maira Sami, Rizwan Ahmed Khan, and Mueen Uddin, "Handwritten optical character recognition (OCR): A comprehensive systematic literature review (SLR)," IEEE Access, vol. 8, pp. 142642-142668, 2020.
- 5) Drobac, Senka, and Krister Lindén, "Optical character recognition with neural networks and post-correction with finite state methods," International Journal on Document Analysis and Recognition (IJDAR), vol. 23, no. 4, pp. 279-295, 2020.
- 6) Yang, Hong-Ming; Zhang, Xu-Yao; Yin, Fei; Sun, Jun; Liu, Cheng-Lin, "[IEEE 2018 16th International Conference on Frontiers in Handwriting Recognition (ICFHR) - Niagara Falls, NY, USA (2018.8.5-2018.8.8)]," 2018 16th International Conference on Frontiers in Handwriting Recognition (ICFHR) - Deep Transfer Mapping for Unsupervised Writer Adaptation, 2018.
- 7) Albahli, Saleh, Marriam Nawaz, Ali Javed, and Aun Irtaza, "An improved faster-RCNN model for handwritten character recognition," Arabian Journal for Science and Engineering, vol. 46, no. 9, pp. 8509-8523, 2021.

- 8) Mainkar, V. Vaibhav, Jyoti A. Katkar, Ajinkya B. Upade, and Poonam R. Pednekar, "Handwritten character recognition to obtain editable text," In 2020 International Conference on Electronics and Sustainable Communication Systems (ICESC), IEEE, pp. 599-602, 2020.
- 9) Hijam, Deena, Sarat Saharia, and Yumnam Nirmal, "Towards a complete character set Meitei Mayek handwritten character recognition," In 2018 Fourth International Conference on Computing Communication Control and Automation (ICCUBEA), IEEE, pp. 1-5, 2018.
- 10) Devi, D. Prabha, et al., "Design and simulation of handwritten recognition system," Materials Today: Proceedings, vol. 45, pp. 626-629, 2021.
- 11) B. M. Vinjit, et al., "A Review on Handwritten Character Recognition Methods and Techniques," 2020 International Conference on Communication and Signal Processing (ICCSP). IEEE, 2020.
- 12) Dey, Raghunah, and Rakesh Chandra Balabantaray, "A reduced feature representation scheme for offline handwritten character recognition," Data engineering and intelligent computing. Springer, Singapore, pp. 629-637, 2021.
- 13) Sahoo, Satyasangram, Prem Kumar, and R. Lakshmi, "Offline handwritten character classification of the same scriptural family languages by using transfers learning techniques," 2020 3rd International Conference on Emerging Technologies in Computer Engineering: Machine Learning and Internet of Things (ICETCE). IEEE, 2020.
- 14) Altwaijry, Najwa, and Isra Al-Turaiki, "Arabic handwriting recognition system using convolutional neural network," Neural Computing and Applications, vol. 33, no. 7, pp. 2249-2261, 2021.
- 15) Carbune, Victor, Pedro Gonnet, Thomas Deselaers, Henry A. Rowley, Alexander Daryin, Marcos Calvo, Li-Lun Wang, Daniel Keysers, Sandro Feuz, and Philippe Gervais, "Fast multi-language LSTM-based online handwriting recognition," International Journal on Document Analysis and Recognition (IJDAR), vol. 23, no. 2, pp. 89-102, 2020.
- 16) Ingle, R. Reeve, Yasuhisa Fujii, Thomas Deselaers, Jonathan Baccash, and Ashok C. Popat, "A scalable handwritten text recognition system," In 2019 International Conference on Document Analysis and Recognition (ICDAR), IEEE, pp. 17-24, 2019.
- 17) S. Kowsalya, P. S. Periasamy, "Recognition of Tamil handwritten character using modified neural network with aid of elephant herding optimization," Multimedia Tools and Applications, 2019.
- 18) Puri, Shalini, and Satya Prakash Singh, "An efficient Devanagari character classification in printed and handwritten documents using SVM," Procedia Computer Science, vol. 152, pp. 111-121, 2019.
- 19) Bora, Mayur Bhargab, Dinthisrang Daimary, Khwairakpam Amitab, and Debdatta Kandar, "Handwritten character recognition from images using CNN-ECOC," Procedia Computer Science, vol. 167, pp. 2403-2409, 2020.
- 20) Deore, Shalaka Prasad, and Albert Pravin, "Devanagari handwritten character recognition using fine-tuned deep convolutional neural network on trivial dataset," Sādhana, vol. 45, no. 1, pp. 1-13, 2020.
- 21) Das, Abhishek, Gyana Ranjan Patra, and Mihir Narayan Mohanty, "LSTM based Odia handwritten numeral recognition," In 2020 international conference on communication and signal processing (ICCSP), IEEE, pp. 0538-0541, 2020.
- 22) Michael, Johannes, Roger Labahn, Tobias Grüning, and Jochen Zöllner, "Evaluating sequence-to-sequence models for handwritten text recognition," In 2019 International Conference on Document Analysis and Recognition (ICDAR),. IEEE, 2019.

- 23) Zhan, Hongjian, Qingqing Wang, and Yue Lu. "Handwritten digit string recognition by combination of residual network and RNN-CTC." In International conference on neural information processing, Springer, Cham, pp. 583-591, 2017.
- 24) Hamdan, Yasir Babiker. "Construction of Statistical SVM based Recognition Model for Handwritten Character Recognition." Journal of Information Technology, vol. 3, no. 02, pp. 92-107, 2021.
- 25) Narayan, Adith, and Raja Muthalagu, "Image Character Recognition using Convolutional Neural Networks," 2021 Seventh International conference on Bio Signals, Images, and Instrumentation (ICBSII). IEEE, 2021.
- 26) Inunganbi, Sanasam, Prakash Choudhary, and Khumanthem Manglem Singh, "Local texture descriptors and projection histogram based handwritten Meitei Mayek character recognition," Multimedia Tools and Applications, vol. 79, no. 3, pp. 2813-2836, 2020.
- 27) Hazra, Abhishek, Prakash Choudhary, Sanasam Inunganbi, and Mainak Adhikari, "Bangla-Meitei Mayek scripts handwritten character recognition using convolutional neural network," Applied Intelligence, vol. 51, no. 4, pp. 2291-2311, 2021.
- 28) Inunganbi, Sanasam, Prakash Choudhary, and Khumanthem Manglem, "Handwritten Meitei Mayek recognition using three-channel convolution neural network of gradients and gray," Computational Intelligence, vol. 37, no. 1, pp. 70-86, 2021.
- 29) Sanasam, Inunganbi, Prakash Choudhary, and Khumanthem Manglem Singh, "Line and word segmentation of handwritten text document by mid-point detection and gap trailing," Multimedia Tools and Applications, vol. 79, no. 41, pp. 30135-30150, 2020.
- 30) Khuman, Yanglem Loijing Khomba, H. Mamata Devi, T. Romen Singh, and N. Ajith Singh, "Graphics Separation System for Printed Document Images," In 2020 International Conference on Computer Communication and Informatics (ICCCI), IEEE, pp. 1-4, 2020.
- 31) Khuman, Yanglem Loijing Khomba, H. Mamata Devi, and N. Ajith Singh, "Entropy-based skew detection and correction for printed meitei/meetei script ocr system," Materials Today: Proceedings, vol. 37, pp. 2666-2669, 2021.
- 32) Khuman, Y. Loijing Khomba, H. Mamata Devi, and Ksh Nareshkumar Singh, "Segmentation of Printed Meitei/Meetei Script Documents," Digital Image Processing, vol. 10, no. 3, pp. 40-44, 2018.
- 33) Khomba Khuman, Yanglem Loijing; Devi, Salam Dickeeta; Devi, H Mamata; Singh, N Ajith; Ponykumar Singh, Chingakham, "A Benchmark Dataset for Printed Meitei/Meetei Script Character Recognition", Mendeley Data, V2, 2022. doi: 10.17632/rw4b2zdk95.2.