

## ENHANCED CHARACTER RECOGNITION MODEL USING BACK PROPAGATION ALGORITHM

Promise Jumbo\*<sup>1</sup>, Umejuru Daniel<sup>2</sup>

<sup>1</sup>Department of Information Technology, University of Port Harcourt, Choba, Nigeria.

<sup>2</sup>Department of Computer Science, University of Port Harcourt, Choba, Nigeria.

Article Received: 18 June 2026, Article Revised: 08 July 2026, Published on: 28 July 2026

\*Corresponding Author: Promise Jumbo

Department of Information Technology, University of Port Harcourt, Choba, Nigeria.

Doi: <https://doi-doi.org/101555/ijrpa.4638>

### ABSTRACT

This research work focuses on the study of recognition of textual characters on natural images using matlab. The research works on joint picture group (jpg) image type. It further utilizes matlab inbuilt processing libraries so characters can be recognized and displayed. The research also employed the principle of Artificial Intelligence on the perception of the environment and to act on the particular image in an environment intelligently. Matlab as a tool and Artificial Intelligence in a neural network was used to improve the artificial computing and experience of science for image optical character recognition. The model can be used to recognize the English characters (A-Z,a-z),numerals(0-9) and special characters(#,\$,%,^,&,\*). Recognition is done by training the image in the matlab IDE using back propagation algorithm, where users input is read as an image when processed to the matlab matrix system and then, the algorithm processes the matrix to know the exact character or number entered by the user and displayed. Our methodology is based on machine learning, especially deep learning principle. The performance of the proposed model shows that efficiency rate of 97% was achieved.

**KEYWORDS:** Images, Character, Recognition, Matlab, AI, Algorithm.

### 1.0 INTRODUCTION

The proposed system is based on textual characters on natural images recognition using matlab. The study works only on joint picture group (jpg) image type. A Matlab graphical user interface is use to encapsulate the steps involved with training an image optical character recognition system. This graphical user interface permits the user to load images, binaries and

segment them, compute and plot features, and save these features for future analysis. Images can be imported into the graphical user interface by clicking on the image menu and selecting open. The proposed system support the JPG image file formats. After opening an image, it binaries and segment them by converting to black and white and segmented them in a different button which will also extract the various features (Szmurlo, 2023). (Nisha Sharma et al, 2023) Image-to-text conversion is also called Optical Character Recognition (OCR), is a crucial technology for Digital Humanities, connecting historical documents written on traditional media (e.g., paper, cloth) to modern, digitized storage. The main domain of the proposed system is to take an image like a photograph of a page of printed or handwritten page of text, and retrieve the original text (Miroslav, 2021). The artificial intelligence in the proposed system is to teach or train the matlab machine to perform this task and scale up the speed and volume that can be handled, so that the billions of pages of historically important written records can be converted from image form to text. The benefit of the research is that the text can be searched and edited with standard computer software, such as Microsoft Word. Furthermore, the text can be processed with Natural Language Processing (NLP) software to study it from the linguistics point of view. The important of optical character recognition is to facilitate access and reduces the cost of storing of massive quantities of data that is used in digital humanities. Artificial Intelligence (AI) is defined as constructing machines or software capable of performing tasks at a level of a competent human (Vijay et al, 2023). Therefore, image-to-text conversion is considered a branch of AI, and makes the proposed system an AI project. Machine learning (ML) has become a feasible approach to many AI problems. The premise of machine language is to build software by training the dataset rather than programming. The components of the software that learn and retain information from the training sets are called Artificial Neural Networks, and are inspired by the biological neural network. The matlab are typically trained using supervised learning. In image-to-text conversion the network may be presented with a few pages of correctly transcribed text, and the network is then capable of transcribing all similar texts. Matlab has a special toolbox, called neural network toolbox which makes the implementation less difficult but the knowledge of theory is needed. Visual information is the most important type of information perceived, processed and interpreted by the human brain; recognition of patterns is automatically done and data are stored. Nosrati, N.(2023), But in the computer world, image processing covers a gamut of scopes such as computerized photography, medical/biological image processing, automatic character recognition, finger print and face recognition, pattern recognition, machine / robot vision ; and it is used for many applications

like industrial applications and security applications on various platforms and infrastructures like distributed systems and clouds (Ronak, K 2022).

## **2.0 Related Works**

Character recognition is the ability of a machine to receive and interpret handwritten input from multiple sources like paper documents, photographs, touch screen devices etc. Recognition of handwritten and machine characters is an emerging area of research and finds extensive applications in banks, offices and industries. The main aim of this project is to design expert system for, “HCR using Matlab” that can effectively recognize a particular character of type format using the matlab network approach. Handwritten character recognition, also called Image-to-text conversion is a crucial technology for Digital Humanities, connecting historical documents written on traditional media like paper and cloth to modern, digitized storage. Sankaran, V. (2022), the main problem of the field sounds deceptively simple: to take an image, e.g., a photograph of a page of printed or handwritten page of text, and retrieve the original text. More precisely, we essentially reproduce the sequence of keystrokes that would result in text semantically equivalent to the content of the original page of text. Of course, a human expert familiar with the writing system used, and capable of typing or other form of data entry, may perform this task. The challenge is to teach a machine to perform this task and scale up the speed and volume that can be handled, so that the billions of pages of historically important written records can be converted from image form to text. The benefits of having the text are numerous. Foremost, the text can be searched and edited with standard computer software, such as Microsoft Word or Emacs (our favorite text editor).

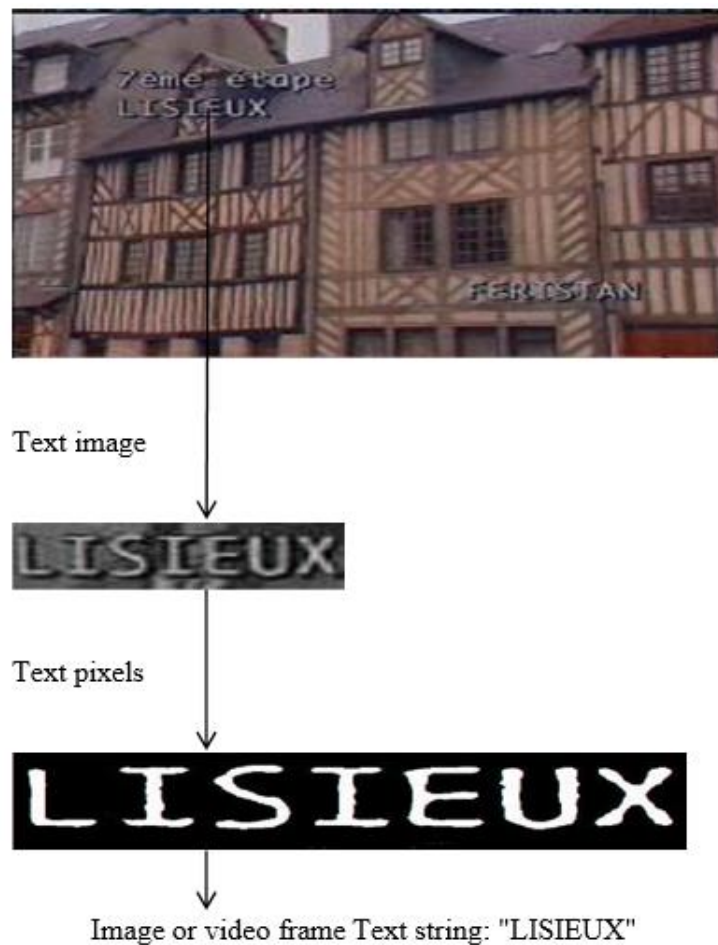
## **2.1 Text Detection and Recognition in Image and Videos**

Text detection and recognition in images and videos is a research area which attempts to develop a computer system with the ability to automatically read from images and videos the text content visually embedded in complex backgrounds. As an example of the general object recognition issues, this computer system, shown in figure 2.1, should answer two typical questions “Where and What”: “where is a text string?” and “what does the text string say” in an image or a video (Wu et al. 2023). In other words, using such a system, text embedded in complex backgrounds can be automatically detected and each character or word can be recognized. The investigation of text detection and recognition in complex background is motivated by cutting edge applications of digital multimedia. Today more and more audio

and visual information is captured, stored, delivered and managed in digital forms. The wide usage of digital media files provokes many new challenges in mobile information acquisition and large multimedia database management. Among the most prominent are:

1. Automatic broadcast annotation: creates a structured, searchable view of archives of the broadcast content;
2. Digital media asset management: archives digital media files for efficient media management;
3. Video editing and cataloging: catalogs video databases on basis of content relevance;
4. Library digitizing: digitizes cover of journals, magazines and various videos using advanced image and video character recognition;
5. Mobile visual sign translation: extracts and translates visual signs or foreign languages for tourist usage, for example, a handheld translator that recognizes and translates Asia signs into English or French.

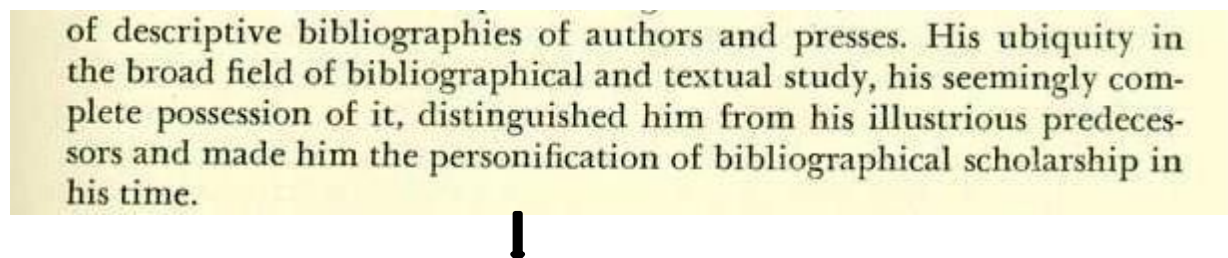
Image or video frame Text string: "LISIEUX"



**Figure 1: Text detection and recognition in images and videos (Zhong, 2023)**

## 2.2 Optical Character Recognition

Optical character recognition (OCR) addresses the problem of reading optically processed characters and has become one of the most successful applications of technology in the field of pattern recognition and artificial intelligence. The start of optical character recognition was motivated by the requirement of high speed and electronic data processing (Paul, 2024). The first optical character recognition equipment installed at the Reader's Digest in 1954 was employed to convert typewritten reports into computer readable punched cards. The matching procedure usually consists of a pre-processing step, a recognition step and an optional post-processing step (Beijing, 1996). In the pre-processing step, the system digitizes a paper-made document into a gray scale image using an optical scanner and converts the gray scale image into a binary image.



**Figure 2: Scanned image of text and its corresponding recognized representation (Szmurlo, 2023) of descriptive bibliographies of authors and presses. His ubiquity in the broad field of bibliographical and textual study, his seemingly complete possession of it, distinguished him from his illustrious predecessors and made him the personification of bibliographical scholarship in his time.**

Some examples of optical character recognition applications are. The most common use optical character recognition is the first item; people often wish to convert text documents to some sort of digital representation.

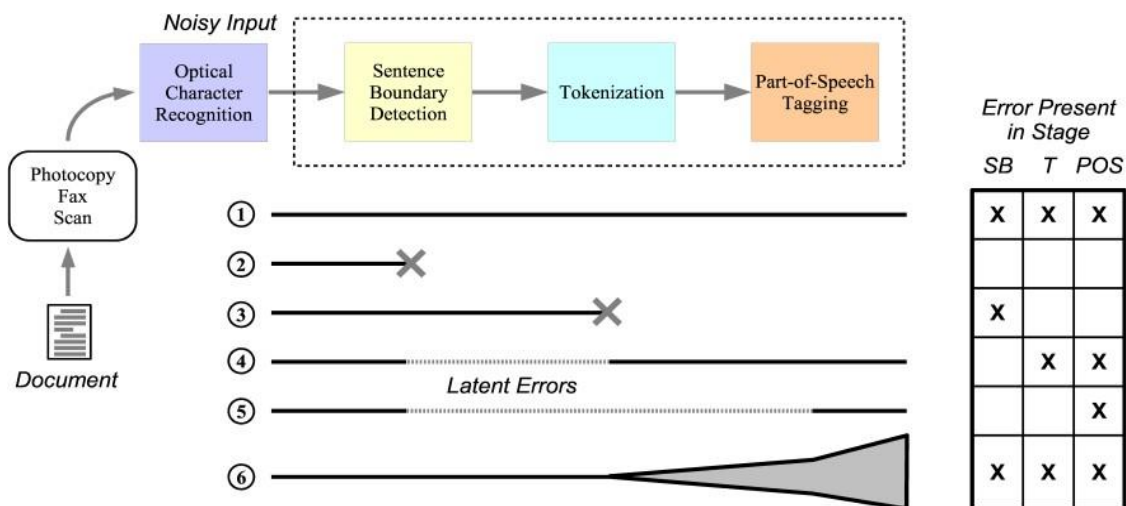
1. People wish to scan in a document and have the text of that document available in a word processor.
2. Recognizing license plate numbers
3. Post Office needs to recognize zip-codes

### The Classification Process

Classification in general for any type of classifier. There are two steps in building a classifier: training and testing. These steps can be broken down further into sub-steps.

## 1. Training

- Pre-processing: Processes the data into a suitable form.
- Feature extraction: Reduce the amount of data by extracting relevant information. Usually results in a vector of scalar values.
- Model Estimation: from the finite set of feature vectors, need to estimate a model, usually statistical for each class of the training data.



**Figure 3: Propagation of OCR mistakes through NLP stages (Miguel, A. 2016)**

A short summary of every stage in addition to its potential problem areas is listed in Table 1. The interactions between mistakes that arise throughout OCR and later stages could be complex. A number of common scenarios are depicted in Figure.2.4. It is straightforward to visualize a solitary mistake proliferating the entire way to the channel and recall an equivalent mistake at all of the afterwards stages in the procedure. Mittrakshi, B. (2012) Mistakes are unavoidable in complex processor image application such as visual quality detection; with the sound induce by these mistakes present a serious face up to do upstream process that efforts to create make utilize of such information.

**Table.1: Processing steps of manuscript.**

| Processing step                         | Intended occupation                                        | Potential difficulties                                                      |
|-----------------------------------------|------------------------------------------------------------|-----------------------------------------------------------------------------|
| <b>OCR</b>                              | Copy contribution bitmap into determined manuscript        | Current OCR is “brittle;” mistakes made early-on propagate to later stages. |
| <b>Sentence boundary reorganization</b> | Break input into sentence-sized units, one per text line.  | Absent or bogus phrase limitations because of OCR mistakes on punctuation.  |
| <b>Tokenization</b>                     | Crack all phrase into word tokens enclosed by blank space. | Absent or bogus tokens because of OCR mistakes on blank                     |



|                                            |                                                                                                 |                                                                                           |
|--------------------------------------------|-------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------|
|                                            |                                                                                                 | space and punctuation.                                                                    |
| <b>Division of language classification</b> | Obtains tokenized manuscript and envoys label to every token representative its part of speech. | Awful PoS marks because of failed tokenization or OCR mistakes that change orthographies. |

### 3.0 METHODOLOGY

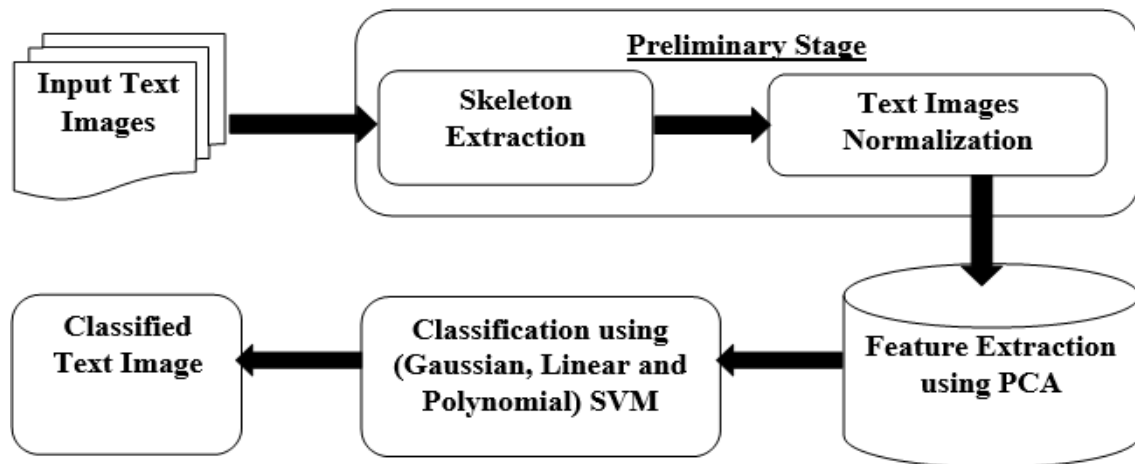


Figure 4: Existing System Architecture (Source: Miroslav, N. 2011).

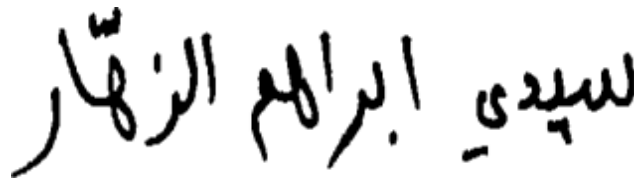
#### 3.1 Description of the Existing System Components

i. **input text image:** the input text image is the image that goes in to the model for conversion from its original image to the require text which can be used by the user. The image is gotten from difference source:

ii. **Preliminary Stage:** In the preliminary stage, text skeleton is first extracted and images of the text are then normalized into uniform size for the purpose of extraction of the universal features of the Arabic text using principal component analysis (PCA). The extracted features are then used to classify the handwritten Arabic text using the SVM classifiers. The preliminary processing stage prepares the text data under consideration for the successive stages. It enhances uniformity of the texts, which is an essential requirement of the recognition system. Preliminary processing is concerned with representation of the text images and normalization processes. At this stage, the skeleton of the word is first extracted by means of the skeletonization-based morphological process so as to eliminate the unnecessary pixels via extraction of the text skeleton at the single-pixel width level. Afterwards, the image of the text of concern is normalized into suitable size for extraction of the global features of the text using the principal component analysis (PCA) technique. A briefing of the two operations making up this stage follows.

iii. **The Skeleton Extraction Process:** In the skeleton extraction process, skeleton of the

input text image is extracted via the skeletonization-based morphological operation method. As such, this process refines text shape and minimizes the size of the data that needs handling for the purpose of feature extraction and recognition. This particular approach was selected to thin the handwritten text because it proved to be having high performance in thinning the handwritten text. An example on skeletons of handwritten Arabic texts that have been extracted using the skeletonization-based morphological method is presented in Figure 3.2 below.



(a) Original Image.



(b) Produced Skeleton.

**Figure 5: Example of (a) Handwritten Arabic Text before thinning and (b) Extracted Skeleton of this Text. (Source: Miroslav, N. 2011).**

**iv. Text Image Normalization:** Normalization of the text images is a very important step in the text recognition process. Because the styles of writing differ from one person to another, the size normalization process is often employed to convert the sizes of the characters or words into a uniform standard size. recognition performance of this system was tested on text images of varying sizes, taking into consideration the smallest and largest image sizes in the relating databases so as to select the best results and compare performance.

**V. Feature Extraction Using Principal Component Analysis:** Feature extraction is the most important process in the recognition systems of the handwritten texts. The eventual goal of feature extraction is to produce efficient representation of the entire text image through set of features. Principal Component Analysis (PCA) is a statistical linear transform technique. The PCA extract and select the relevant features of handwritten text as a global statistical text feature extraction technique for the extracted features to be classified by the SVM classifiers. The PCA is commonly employed as feature extraction method so as to reduce dimensions of



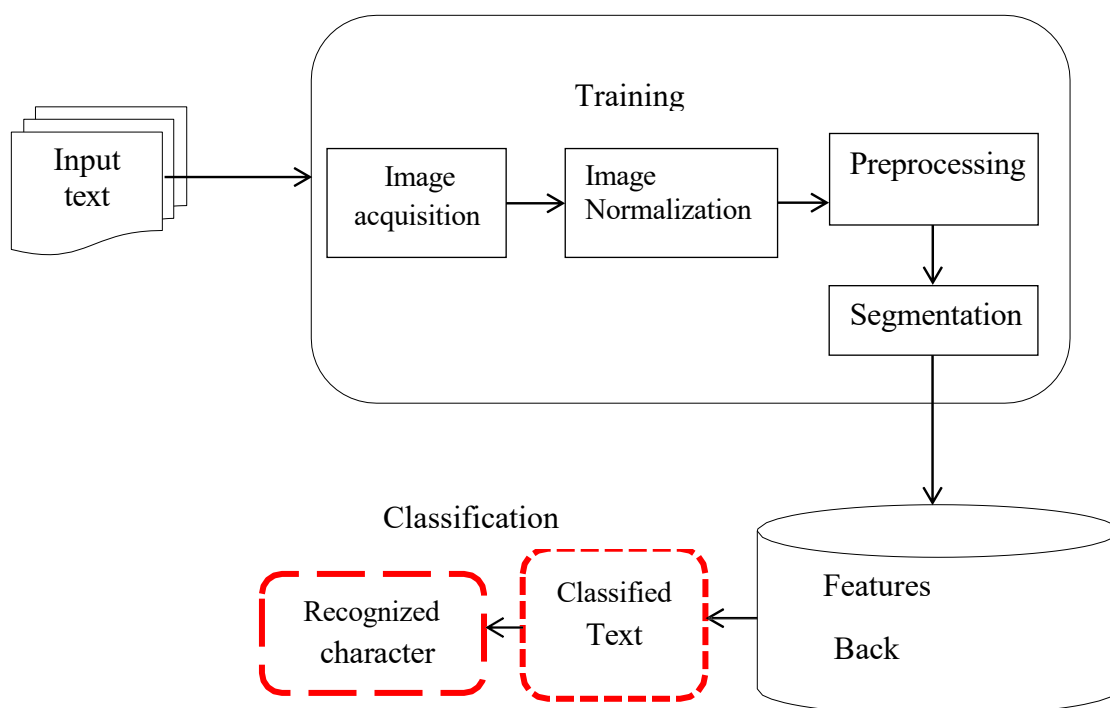
images to manageable sizes. PCA begins by calculating the mean of the data matrix. Then, it computes the covariance of the data. Thereafter, the Eigenvalues and Eigenvectors are estimated. The PCA aims at finding the space that represents direction of the maximal variance of the data under consideration. It defines low-dimensional space, or a PCA space (W), that can be employed to transform the data ( $X = \{x_1, x_2, \dots, x_n\}$ , where n is number of observations or samples and  $x_i$  is the  $i$ th observation, sample, or pattern) from high-dimensional space into low-dimensional space.

**vi. Classification using SVM Classifiers:** The SVM is a relatively modern classifier that employs kernels to find the optimum decision boundary and, then, separate between the potential classes in the high-dimensional feature spaces. The fundamental form of the linear SVM classifier attempts to discovery the optimum hyper plane that separate the best set of samples that belong to differing classes. In classification component, the multi-class polynomial, Gaussian, and linear SVM classifiers were used for classification of images of handwritten text by using the sequent SVM kernel functions:

The linear function:  $K(x, y) = (K(x_i, x_j) = (x_i x_j))$

The polynomial function:  $K(x, y) = (K(x_i, x_j) = (\gamma x_i x_j + \text{coef}))$  The Gaussian function:  $K(x, y) = \exp(-\gamma \|x_i - x_j\|^2)$

**vii. classified Text Image:** The classified text image is the require text that come from the classification component. The user can then use the classified text for it personal function.



**Figure 6: Proposed System Architecture.**

### 3.2 Descriptions of the Proposed System Components

**System Training Module:** This module can be accessed by both the administrator and the end-user. Before converting the printed documents in to editable and searchable documents, the first and the mandatory step is providing training to the system. Here training in the sense the font followed in the scanned document should be identified by the user. Then the user types all the characters that are required for recognition from the scanned document as an image file. This image file should be provided as an input during the training process. The user then clicks the train button provided in the recognition module. Then the training gets completed. Thus the system gets familiar with the new font.

i. **The input text image:** these are the dataset images that goes in to the model for conversion from its original dataset image to the require text which can be used by the user. The image is gotten from difference source which include: optical scanning of images of various handwritten documents, by clicking photographs through cameras, by coping of text image from textbook.

ii. **Image Acquisition:** The proposed character recognition system starts with image acquisition process that takes an input of a digital image by using a digital camera or scanner. Specific format such as JPEG, PNG are taken as the input image. This image is used for testing and training the back propagated algorithm.

iii. **Image Normalization:** the image and the format in the acquisition component go into the image normalization. The style of writing the images differ from one person to another and the size of the image is also differ. Size normalization process is use to convert the sizes of the characters or words into a uniform standard size.

vi. **Preprocessing:** Many times while inputting the image to the system for handwriting recognition process the image rotates itself, many times it gets ruptured with noise, blurring of image takes place, size of the text in image can be smaller or larger in different cases and many problems come across in this stage. So in preprocessing stage, basically the cleaning of input image is done so that the resulting image after preprocessing can become suitable for further stages of recognition process. Following steps are performed for carrying preprocessing in the proposed system are:

i. Initially, the colored image is converted into gray scale image using in-built function “rgb2gray”.

ii. Then filtration is done in which the noise is removed from the input image. Filter is used to filters binary image that has been corrupted with constant power. It is a low pass filter.

iii. After this, the input image is converted into “BW”. The output BW image replace all

pixels in image with a value greater than a specified threshold value with value 1 (white) and all other pixels are replaced with 0 (black).

iv. In the end, normalization is done in which image text is normalized to a specified range.

**v. Segmentation:** Segmentation is done to make the separation between characters in an image. There are different ways in which characters interfere: overlapping, touching, connected and intersecting pairs. For segmenting handwritten digits, lines-words-digit segmentation method is use.

i. handwritten text is first segmented into lines, such that text in an image can be written in multiple lines.

ii. Initially each line having text is counted.

iii. After that, each line is evaluated for further segmentation process.

vi. In each line, words are segmented. This is called word segmentation.

v. After that, character segmentation takes place in which each word is segmented in single character or letter. This happens in a sequential way that is initially lines are segmented, then first line is evaluated in which word segmentation is done and in that line character segmentation is done.

vi. After the complete evaluation of first line second line is evaluated for segmentation process.

**vi. Feature Extraction and Digits Recognition using:** the extraction improves the recognition rate of the system and reduces the misclassification of the handwritten text while recognition process. Feature extraction is use for finding the set of unique features which can define the shape of a character precisely. In this proposed work, classical technique Template matching is used by the matlab machine. This technique is also called as Correlation template matching technique. Following steps are performed:

i. Templates of digits are created for training the system such that system remembers those templates.

ii. These templates are stored in the database for future use. When the recognition process starts the input digits for testing are matched with those pre-defined templates in the database.

iii. The testing image which matches the most with the pre-defined templates in the database are classified in one class. Recognition process is done.

**vii. Classified Text:** the classified text component contain all the extracted feature from the image. These texts can be used by the user for any function. The user can copy the texts from

the stored database and place in a require environment for reuse.

### 3.3 Algorithm of the Proposed System

The step used in the back propagation algorithm with a two dimension matrix includes: Step

1: Input images of written text

Step 2: Output the classified Word

Step 3: Read the image of the handwritten text

Step 4: prepare the data of the text image using the following two preliminary steps:

i. Extract the text skeleton and feature text using the two-dimensional (2D) text image

ii. Normalize the text image size to a suitable size

Step 5: The two-dimensional (2D) text image is transformed into mono-dimensional vector by concatenating every column (or row) in the 2D matrix in order to create long vector. If we have an M vector of the size N that represents sample image, then  $X_i = [p_1 \dots p_n]^T$   $i = 1, \dots, N$

Where  $X_i$  is a vector,  $p_x$  is pixel value  $X_i$ , and T is transpose of the vector set. Step 6: Input x: Set the corresponding activation for the input layer.

Step7: Feedforward: For each  $l=2,3,\dots,L$  compute  $z_l = w_l a_{l-1}$  Step8: Output error  $\delta_L$  Compute the vector.

Step9: Back propagate the error: For each  $l=L-1, L-2,\dots,2$  compute  $\delta_l = w_{l+1} \delta_{l+1}$  Step10: Output the gradient given.

### 4.0 DISCUSSION OF RESULT

Other extraction done by the back propagation algorithm is to extract each word by squaring it to identify the word. It use object annotation to write the text on top of the selected square word. After extraction, it then displays the word line by line there after display the text based on the original written image using ocr as the result.

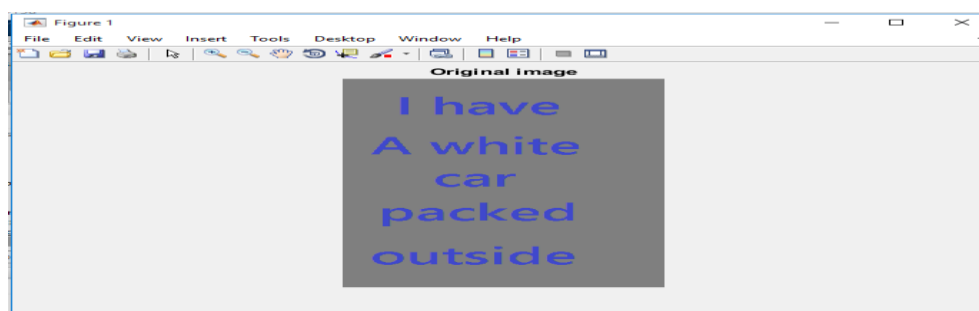
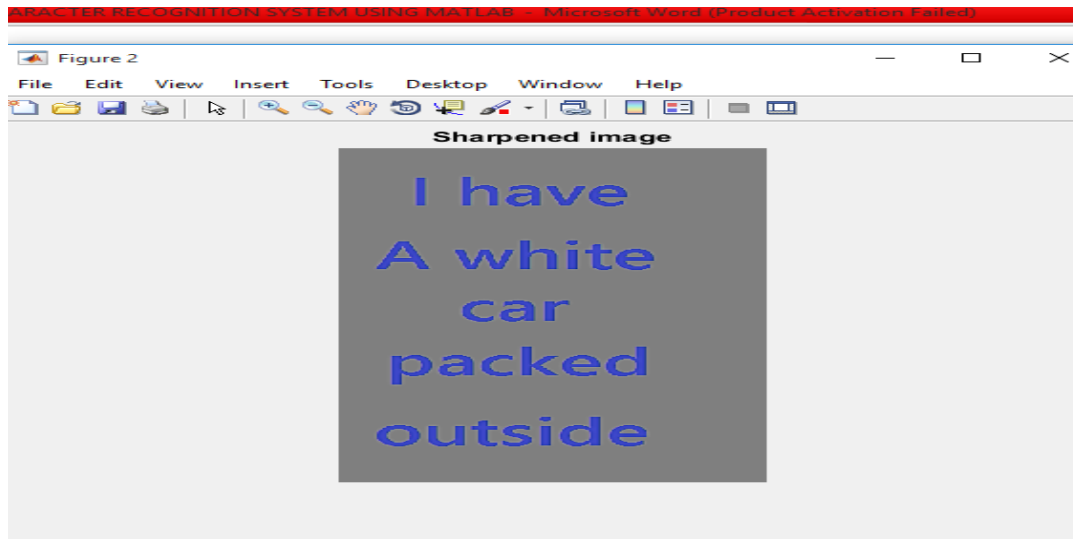


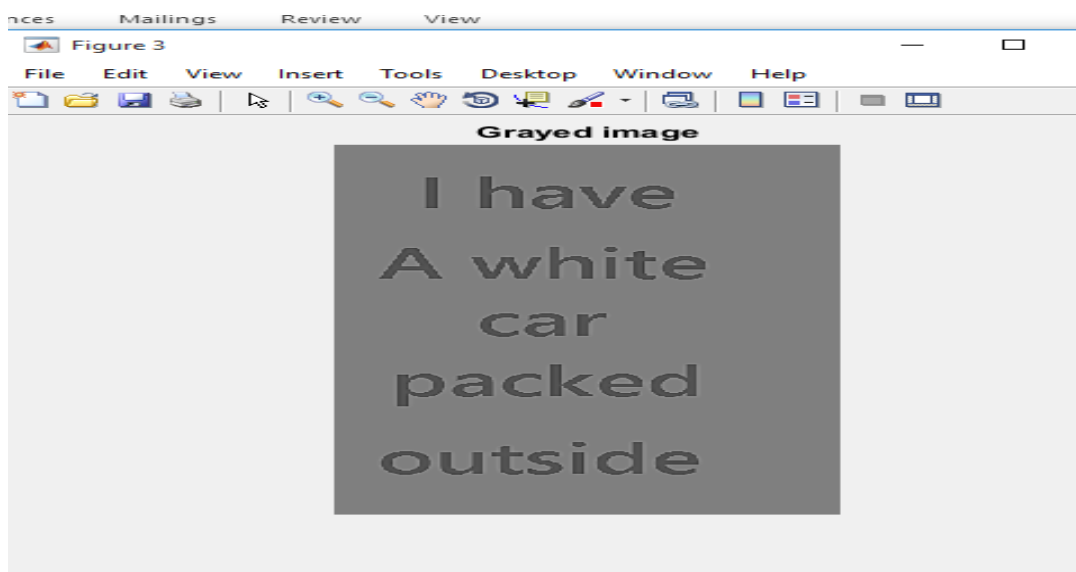
Figure 7: The Original Image.

Figure 7 shows the original image, this image is the input text image that needs to convert to the desire text. The original image is gotten from the car .png (portable network graphic) that is inputted into the matlab machine. This image will be extracted in difference form using the back propagated algorithm.



**Figure 8: Sharpened Image.**

Figure 8 shows the sharpened image, the sharpened image is use to remove noise and distortion from the original image. Noise in the image is the dot found in the image while distortions are the blurred that make images unclear. The function of the sharpened image is to eliminate all these malfunction element in the image.



**Figure 9: Grayed Image.**

Figure 9 depicts the grayed image, the grayed image is use to make the background of the image light and to make the text dark or bold. This function „rgb2gray“ converts the true color image to a grayscale. Intensity image, this is to remove hug illumination an saturation, that may prevent or distort Optical Character Recognition, OCR  $\text{grayImage} = \text{rgb2gray}(\text{imgSharpened})$ .



**Figure 10: No noise Image.**

Figure 10 shows the noiseless image, the no noise image use an adaptive noise removal and low-pass filter to remove the noise in the image. It filters the image which its true nature has been removed by a constant power adaptive noise.

**Table 2: Performance Evaluation of the Proposed System.**

| S/N | PARAMETER    | NUMBER OF TEXT IN IMAGE | COLOURE IMAGE TEXT | EFFICIENCY (%) |
|-----|--------------|-------------------------|--------------------|----------------|
| 1   | Car.jpg      | 6                       | Black              | 95             |
| 2   | Ai.jpg       | 3                       | Yellow             | 67             |
| 3   | Book.jpg     | 5                       | Mix-colour         | 75             |
| 4   | Bar.jpg      | 7                       | Yellow             | 65             |
| 5   | Goat.jpg     | 5                       | White              | 90             |
| 6   | Downing.jpg  | 4                       | Red                | 55             |
| 7   | Hba.jpg      | 7                       | mix-colour         | 50             |
| 8   | Princess.jpg | 1                       | Black              | 95             |
| 9   | School.jpg   | 5                       | Blue               | 85             |

Table 2 contains the parameters use for the proposed system. The evaluation of the efficiency of each dataset parameter is calculated based on the number of text in the input image and the background color of each image. The time taken to scanned and extract individual image text



depend on the number of text and the background color. The efficiency of each parameter is recorded in percentage. For verifying performance of the proposed system, the recognition results of the proposed system were compared with those of existing recognition system with the same dataset parameter.



**Figure 11: Performance Evaluation Chart of the Proposed Systems.**

A graph of efficiency against parameter is plotted in figure 11 from the graph the proposed produced better classification accuracies than other existing systems when applying the aforementioned dataset normalized image sizes it will open and display different interface, then change the directory by clicking on the browse for folder. Then click on desktop and click on character recognition, then click on select folder. In the directory you observe many image extensions, all the image extension have been taking from different source. Then double click on update.txt, Loads the image from the working directory into `colorImage = imread('goat.png');` each image can be change in .png extension inside the bracket. Then copy the directory and paste in the command window, then from the command window, execution of extraction of the image begin, at the end of extraction it first display the words on the image before display the final text that can be used by the user.

## 5.0 CONCLUSION

In this proposed research, an approach for recognizing handwritten digit is described. The recognition accuracy of the implementation is promising. In this proposed work, preprocessing work is done in which normalization, filtrations are performed using Weiner

filter so that input image is noise free and become appropriate in further steps for recognition. This work also used digit segmentation method which can handle a larger variety of touching and overlapped digits, which occurs fairly often in images obtained from handwritten material. Finally correlation matrix is used for feature extraction step and features are extracted, and used for number recognition process. The performance of the proposed model shows that efficiency rate of 97% is achieved.

## REFERENCES

1. Acharjya D. (2023), A Survey on Big Data Analytics: Challenges, Open Research Issues and
2. Tools, International Journal of Advanced Computer Science and Applications, (IJACSA) 7(2), 511 – 518
3. Anita Thengade, RuchaDondal, Genetic Algorithm – Survey Paper, MPGI National Multi Conference 2012 (MPGINMC - 2020) 7-8 April, 2020 “Recent Trends in Computing” Proceedings published by International Journal of Computer Applications. 0975- 8887.
4. Alom, P. (2020) “Handwritten Bangla Digit Recognition Using Deep Learning,” 6( 7): 990–997, . Beniwal et al, (2024), “Survey on Handwritten Digit Recognition using Machine Learning,” Int. J. Comput. Sci. Eng., 06(05): 96–100.
5. Chen, Z. (2020), “Handwritten digits recognition,” Proc. 2009 Int. Conf. Image Process. Comput. Vision, Pattern Recognition, IPCV, 5(2): 690–694.
6. Cheedella, G. (2020), “Critique of Various Algorithms for Handwritten Digit Recognition Using Azure ML Studio,” Glob. J. Comput. Sci. Technol., 20(1): 1–5.
7. Christina P. (2021), A Parallel and Scalable Processor for JSON Data, Industrial and Applications Paper, 11(13), 47 – 51
8. Majumder, S. (2023) “Handwritten Digit Recognition by Elastic Matching,” J. Comput., 4(4): 1067–1074.
9. Miguel, 2023). A Genetic Algorithm-Based Clustering Approach for Database Partitioning. IEEE Transactions on Systems, Man, and Cybernetics, 32(3), 215-230.
10. Miroslav (2021) Optical Character Recognition of Printed Persian/Arabic Documents. Electronic Theses and Dissertations, 21(7): 231-432.
11. Mittrakshi, B. (2022) “Sindhi Handwritten-Digits Recognition Using Machine Learning Techniques Sindhi Handwritten-Digits Recognition Using Machine Learning Techniques,”
12. Nisha Sharma et al (2023),, “A Study of Moment Based Features on Handwritten Digit

- Int. J. Comput. Sci. Netw. Secur., 19(5): 195–202.
13. Norati N. (2023), proposed using a statistical Gaussian mixture model, and investigated simple training data preprocessing and focused on improving recognition rates using committees of neural networks. 2(7): 67-103.
  14. Ronak K (2022) variance of edge orientation the method take care of characteristics coming from parallel character edges 2(7):1030–1033
  15. Rokus Arnold (2020) Detecting Text Lines in Handwritten Documents. In 18<sup>th</sup> International Conference on Pattern Recognition (ICPR'06), 23(4): 130–143.
  16. Sankaran, C. (2022) Recognition of printed Devanagari text using BLSTM Neural Network. In Proceedings of the 21st International Conference on Pattern Recognition (ICPR2012), 322–325.
  17. Shamim, M. (2021), “Handwritten digit recognition using machine learning algorithms,” Indones. J. Sci. Technol. 3(1): 29–39, 2018.
  18. Shyla afrogee et al (2022) A Long Short-Term Memory Recurrent Neural Network Framework for Network, International Journal of Computer Applications (IJCA), 7(34), 22-27
  19. Smith et al (2025) Genetic algorithms-based approach to database vertical partitioning. Journal of Intelligent Information Systems, 26 (2), 167-183.
  20. Suen, and Bloch, (2021), “A trainable feature extractor for handwritten digit recognition,” Pattern Recognit., 40(6): 1816–1824
  21. Szmurlo, (2023) Robust Recognition of Degraded Documents Using Character N-Grams. In Proceedings of the 10th IAPR International Workshop on Document Analysis Systems, DAS 9(12): 130–134.
  22. Vijay et al (2023) “Handwritten digit recognition: A neural network demo,” Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics), LNCS, 32(21): 762–771.
  23. Yousaf, “Handwritten Digit Recognition Using Image Processing and Neural Networks,” Lect. Notes Eng. Comput. Sci. 2165(1): 648–651.
  24. Wu et al. (2023). “Machine Learning for Handwriting Recognition,” 43(4523): 93–101. Recognition,” Appl. Comput. Intell. Soft Comput., 1–17.