

Handwritten Character Recognition and Prediction of Age, Gender and Handedness Using Machine Learning



Inspiring Excellence

School of Engineering & Computer Science

**Department of Computer Science and Engineering,
BRAC University**

Supervisor: Dr. Jia Uddin

Submitted by:

Md. Al Emran - 14301030

S. M. Naief -14101122

Md. Shimul Hossain -14301027

DECLARATION

We, hereby declare that this thesis titled “**Handwritten Character Recognition and Prediction of Age, Gender and Handedness Using Machine Learning**” is based on the results found by ourselves. Materials of work found by other researcher are mentioned by reference. This Thesis, neither in whole or in part, has been previously submitted for any degree.

Signature of Supervisor

Dr. Jia Uddin
Assistant Professor,
Department of Computer Science & Engineering
BRAC University

Signature of Author

Md. Al Emran

S. M. Naief

Md. Shimul Hossain

ACKNOWLEDGEMENTS

We would like to thank almighty Allah, the most merciful, beneficent and most gracious, who always guided us to the right path and helped us complete our thesis. There are many people we would like to acknowledge, for their support during the journey of preparing for our thesis.

We thank our thesis supervisor, Dr. Jia Uddin for his constant support in completion of our thesis. His guidance and encouragement for work has helped us learn new things as well as gave us extraordinary experiences throughout the work. We thank all the faculty staffs of Computer Science and Engineering department, BRAC University for guiding us in the whole study period at BRAC University.

We express our sincere gratefulness to our beloved parents, for the overwhelming love and care they bestow upon us. Our deep love and thanks are due to our brothers and sisters. Finally, we would like to thank our friends who helped us directly or indirectly to complete our thesis.

ABSTRACT

Handwritten character recognition and prediction of age, gender & handedness from handwritten documents offers an interesting research problem for researchers as few research carried out on this field. The aim of this research is to investigate machine learning classification algorithm that is used to recognize different writer's attributes and their handwritten characters. Predicting writer's identity and recognizing handwritten characters based on mainly three steps: segmentation, feature extraction and classification. In the segmentation step we used edge detection technique for segmenting dataset images using fuzzy logic. Feature extraction methods are described to take decision category of our writers and their handwritings. For feature extraction we used mRMR for feature selection, tortuosity, direction, curvatures and chain code for feature extraction and PCA for dimension reduction. In the final step, we used KNN, SVM and RFC for classification of writer attributes and recognizing handwritten characters. Classification accuracy on QUWI dataset were 89.41% for recognizing handwritten character, 88.28% for age range prediction, 75.90% for gender prediction and 75.11% for handedness prediction for each writer. We have used these classification algorithms to bring out the maximum accuracy rate for predicting age, gender & handedness.

Keywords: SVM, RFC, KNN, PCA, mRMR, Handwriting Recognition, Chain Code, Tortuosity, Direction, Curvatures, Dimension Reduction, Edge Detection.

DECLARATION	ii
ACKNOWLEDGEMENTS.....	iii
ABSTRACT.....	iv
CONTENTS	v
LIST OF FIGURES.....	vii
LIST OF TABLES.....	viii
LIST OF ATTRIBUTES.....	ix

CHAPTER 1: INTRODUCTION

1.1 General.....	1
1.2 Motivation.....	2
1.3 Thesis Orientation.....	2

CHAPTER 2: LITERATURE REVIEW..... 3

CHAPTER 3: Foundation of Handwriting Recognition & Prediction of Age, Gender and Handedness

3.1 Handwriting Recognition.....	5
3.2 Off-line Recognition.....	5
3.3 Problem Domain Reduction Techniques	5
3.4 Character Extraction.....	5
3.5 Character recognition.....	6
3.6 Neural Networks.....	6
3.7 Feature Extraction.....	6
3.8 On-line Recognition.....	7
3.9 General Process.....	7

CHAPTER 4: Proposed Model for Handwriting Recognition & Prediction of Age, Gender and Handedness

4.1 Data Set Analysis	8
4.2 Proposed Model.....	9
4.2.1 Image Acquisition.....	9
4.2.2 Image Preprocessing.....	10
4.2.3 Image Segmentation using Edge Detection Technique.....	10
4.2.4 Feature Extraction.....	12
4.2.5 Dimension Reduction using PCA.....	16
4.3 Classification.....	20
 CHAPTER 5: EXPERIMENTAL RESULT.....	 23
CHAPTER 6: CONCLUSIONS AND FUTURE WORKS.....	29
CHAPTER 7: REFERENCES.....	31

List of Figures

Figure 4.1: Sample Pages of QUWI Dataset	8
Figure 4.2: Proposed model.....	10
Figure 4.3: Edge Detection Technique.....	12
Figure 4.4: Tortuosity Feature	13
Figure 4.5: Direction feature	14
Figure 4.6: Chain Codes.....	15
Figure 4.7: Curvatures.....	15
Figure 4.8: Extracted Curvature Value. Left: male; Right: female.....	16
Figure 5.1: Performance analysis of every feature for age, gender & Handedness	28
Figure 5.2: Performance analysis of every feature for handwritten character Recognition	29

List of Tables

Table 5.1: Classification accuracy for age.....	24
Table 5.2: Classification accuracy for gender.....	25
Table 5.3: Classification accuracy for handedness.....	26
Table 5.4: Classification accuracy for handwritten character recognition	27

List of Attributes

KNN = K-Nearest Neighbor

SVM = Support Vector Machine

RFC = Random Forest Classifier

QUWI= Qatar University Writer identification

PCA = Principal Component Analysis

mRMR= Minimum Redundancy Maximum Relevance

CHAPTER 1

INTRODUCTION

1.1 General

Handwritten character recognition is the ability of a computer to read and interpret handwritten characters from various sources such as documents, images, touch-screens and various gadgets. Handwritings can be used to identify the author's identity such as age, gender, handedness and ethnicity and also come in a great advantage in signature verification, postal-address interpretation, bank-check processing, etc. A person's handwriting is like a fingerprint: one can copy others handwriting but can never write it in an identical way. This refers that not two individuals have exactly the same handwriting, and thus, analysis of this feature becomes more significant in various sectors like forensic biometrics, automotive industry, business field, service personnel, education, bank-check processing, signature verification, postal-address interpretation, demography-based writer identification etc. Forensic department of police hire handwriting experts with high costs to identify criminals from examining and analyzing their writings. This manual procedure is both time-consuming and expensive, and thus, to solve the issue, artificial intelligence can come into use to provide an accurate outcome in a shorter time. Moreover, handwriting recognition also plays a significant role in authentication of documents and genuineness of historical manuscripts. Both neuroscientists and psychologists agree on the correlation between handwriting and various demographic attributes of writers.

Sarah Hamid and Kate Miriam Loewenthal [1] conducted a research work where people were asked to predict the gender of a writer from a given handwriting sample . A classification rate of 68% was similar to those reported in studies of handwriting in English [1]. Psychology can also be benefited from research on handwriting style since it could be possible to correlate handwriting and personality attributes of the writer.

1.2 Motivation

The objective of this research is to predict age, gender & handedness of writer from handwritings for application in various sectors like healthcare, automotive industry, business field service personnel, education, demography-based writer recognition, signature verification,

postal-address interpretation, bank-check processing etc. In Bangladesh and other countries, kidnapping is a common case. Very often the kidnappers leave only a handwritten note for clue. Another common issue is claiming one's intellectual property by using wrong information. Moreover, other criminal activities found where handwriting recognition plays an important role. To solve these various issues, we need experts to identify the writers of various documents. It takes time and the accuracy rate is much lower. So, using machine learning algorithms we can reduce the time and increase the accuracy rate for handwriting recognition. We believe, our proposed model will provide more accurate result in the way to predict age, gender and handedness from handwriting.

1.3 Thesis Orientation

The rest of the thesis is organized as follows:

- Chapter 2: describes the related work to our research.
- Chapter 3: describes the foundation knowledge of handwriting recognition
- Chapter 4: narrates the proposed model of our research.
- Chapter 5: explains the results found in our research
- Chapter 6: wraps up this thesis with future research possibilities for out thesis.
- Chapter 7: references that we used to do research.

Chapter 2

LITERATURE REVIEW

In this chapter, we would discuss the theoretical and methodological approaches that were established related to our work.

Anuj Dutt [2] in his paper showed that using Deep Learning techniques, he was able to get a very high amount of accuracy. By using the convolutional Neural Network with Keras and Theano as backend, he was getting an accuracy of 98.72%. Moreover, implementation of CNN using Tensorflow gives an even better result of 99.70%. Though the complexity of the process and codes seems more as compared to normal Machine Learning algorithms but the accuracy he got is more appreciable [2].

In another paper, S. Al-ma'adeed and A. Hassaine[3] in their paper proposed a method for the classification of age, gender and nationality through handwritings. Using Random Forest Classifier and Kernel Discriminant analysis using spectral regression for classification. They built the Random Forest Classifiers for the cases of age, gender and nationality using R random forest library. They describe the QUNI handwriting database for their experiments. A random classification would predict approximately 50%, as this is a two-classification for gender classification. A random classification would therefore predict approximately 14% for Age classification. Then, a random classification would only predict approximately 12% for nationality prediction [3].

In their proposed system, Abdul Hamid and Amin Sjarif [4] used three classification algorithms to recognition which are Support Vector Machine (SVM), K-Nearest Neighbor (KNN) and Multilayer Perceptron Neural Network (MLP). Among these algorithms, SVM and KNN correctly predict the dataset but MLP Neural network couldn't predict the number 9 . They also showed the reason behind this case. KNN and SVM predict directly from feature extraction while MLP is a non-linear function. So, it is more suitable for learning non-linear models. Moreover, MLP with hidden layers have non-convex loss function where there exists more than one local minimum. Different random weight initializations can lead to different validation accuracy. But using Convolutional Neural Networks with Keras can help to improve it [4].

Angel Morera, Angel Sanchez, Jose Francisco Velez and Ana Belen Moreno [5] in their paper proposed a method on the application of deep neural networks to several automatic demographic classification problems based on handwriting . They have addressed three problems which are gender, handedness and the combined “gender-and-handedness” classification. They have used Convolutional Neural Network (CNN) for their classification. They used two databases; Offline IAM and KHATT, respectively. For gender classification on IAM database, overall accuracy is 80.72% and for handedness classification, the accuracy is 90.70% and for the combined “gender-an-handedness” classification, average accuracy is 83.19%. They also used KHATT database for gender, handedness and combined gender-and-handedness classification where overall accuracy is 68.90%, 70.91%, 70.84%. Finally, the comparison of their results is that their solution produced the best accuracy results for the gender classification problem on both tested handwriting database [5].

A method to implement classification algorithm to recognize handwritten digits is proposed by Gaurav Jain and Jason KO [6] in their paper. Their goal was to implement a pattern classification method to recognize the handwritten digits provided in MNIST database of images of handwritten digits (0-9). They used K-Nearest Neighbor classification which is one of the simplest classifications to implement. For their project, they chose K to be 1, so they have implemented a 1-Nearest Neighbor classifier. They showed that Convolutional Neural Networks (CNN) is a special kind of multi-layer neural networks used to recognize visual patterns with extreme variability. Then they able to obtain a maximum accuracy of 78.4% [6].

K.Venkata Reddy, D.Rajeswara Rao, U.Ankaiah and K.Rajesh [23] in their paper proposed a method to recognize the handwritten digits and characters by using the multilayer perception artificial neural network. In this MLP network is use the back propagation algorithm to train and test the data. They worked with back propagation algorithm consists of two phases. First phase is forward phase and second phase is backward phase. First phase is the phase where the activations propagate from input layer to the output layer. And the second one is the phase where the error between the observed actual value in the output layer are propagated backwards so it can modify the weights and bias values. They used Artificial neural network is used to recognize ten different handwritten digits. Finally, they presented a new system for handwritten text recognition based on an improved artificial neural network.

Chapter 3

Foundation of Handwriting Recognition & Prediction of Age, Gender and Handedness

3.1 Handwriting Recognition

Handwriting recognition principally entails the character recognition. However, a complete recognition system handles formatting, performs correct segmentations into characters and finds the most plausible words.

3.2 Off-line Recognition

Off-line handwriting recognition includes the programmed transformation of content in an image into letter codes which uses in pc and text-processing application. The information acquired by this frame is viewed as a static representation of handwriting. It is difficult to process because different people have different handwriting styles. What more, as of today, OCR engines fundamentally centered around machine printed content and ICR for hand “printed” content.

3.3 Problem Domain Reduction Techniques

Narrowing the problem domain frequently helps incrementing the accuracy of handwriting recognition systems. A shape field for any zip code would contain only the characters 0-9. However, this fact would decrease the quantity of possible identifications.

Primary Techniques-

- Specifying specific character ranges.
- Utilization of specialized forms.

3.4 Character Extraction

Offline character recognition regularly includes examining a shape or archive composed at some point previously. This implies the individual characters combined in the filtered picture should be removed. Tools exist that are fit for playing out this step [7]. In any case there are a few regular defects in this progression. The most well-known is when character that is associated is returned as a single-sub image containing both characters. This causes a major problem in the recognition stage [8].

3.5 Character Recognition

After the extraction of individual character happens, a recognition engine is utilized to distinguish the comparing PC character. A few diverse recognition methods are right now accessible [8].

3.6 Neural networks

Neural network recognizers gain from an underlying image preparing set. The prepared system at that point makes the character distinguishing pieces of proof. Each neural network exceptionally takes in the properties that separate preparing image. It at that point searches for comparable properties in the objective image to be distinguished. Neural network are quick to set up; nonetheless they can be wrong on the off chance that they learn properties that are not critical in the objective information [8].

3.7 Feature Extraction

Feature extraction works in a similar fashion to neural network recognizers. However, programmers must manually determine the properties they feel are important. Some example properties might be:

- Aspect ratio
- Percent of pixels above horizontal half point
- Percent of pixels to right of vertical half point
- Number of strokes
- Average distance from image center
- Is reflected by y axis.
- Is reflected by x axis.

This approach gives the recognizer more control over the properties used in identification. Yet any system using this approach requires substantially more development time than a neural network because the properties are not learned automatically [8].

3.8 On-line Recognition

Online handwriting recognition includes the programmed change of content as it is composed on an extraordinary digitizer or PDA, where a sensor gets the pen-tip developments and additionally pen-up/pen-down exchanging. This sort of information is known as advanced ink and can be viewed as a computerized portrayal of handwriting. The acquired flag is changed over into letter codes which are usable inside PC and content handling applications [8].

The elements of an on-line handwriting recognition interface typically include [8]:

- A pen or stylus for the user to write with
- A touch sensitive surface which may be integrated with, or adjacent to, an output display.
- A software application which interprets the movements of the resulting strokes into digital text.

3.9 General Process

The process of on-line handwriting recognition can be broken down into a few general steps

- Image acquisition
- Image preprocessing
- Image segmentation
- Feature extraction
- Feature selection and classification

The reason for preprocessing is to dispose of insignificant data in the information that can contrarily influence the recognition [9]. This concerns speed and accuracy. Preprocessing usually consists of binarization, normalization, sampling, smoothing and denoising [10]. The second step is feature extraction. Out of the at least two-dimensional vector field received from the preprocessing algorithms, higher-dimensional data is extracted. The purpose for this progression is to feature virtual data for the recognition model. This information may incorporate data like pen pressure, velocity or the altars of composing course. The last big step is classification. In this progression different models are utilized to outline removed features to various classes and therefore recognizing the characters or words the features represent.

Chapter 4

Proposed Model for Handwriting Recognition & Prediction of Age, Gender and Handedness

4.1 Data Set Analysis

We have used QUWI dataset [12] for handwritten character recognition and prediction of age, gender & handedness. QUWI dataset contains 1017 writers for which the genders are provided [12]. A total of 1017 writers produced 4 handwritten documents. The QUWI dataset is differing. The novelty and genuine preferred standpoint of the dataset are the assorted variety of writers, of dialects and all criteria. The examination of the dataset demonstrates that 306 of the volunteers are Qataris, approximately 190 are Lebanese, 101 are Palestinians, 104 are Egyptians, and 68 are Jordanians [12]. There are too numerous Sudanese, Yemenis, Syrians, Iranians, Iraqis and Saudis people spoke to in the dataset. The variety in the instructive levels of the members is moreover interesting: the dataset includes not just exceptionally taught yet in addition less instructed individuals. In reality, the volunteers included primary school students, secondary school students and university students and also workers, employees, engineers, specialists, professors, accountants and furthermore secretaries in modern, administrative and academic environments. The volunteers additionally varied in age. There are volunteers younger than 12 years, teenagers, and adults and older than 40 years [12]. The dataset was 52% written by females (530 authors) and 48% written by males (487 authors)[12]. 953 of the volunteers (93.7%) are right handed though 64 volunteers (6.3%) are left handed. These factors (nationality, age go, handiness, instructive level, and sex) are critical elements of our dataset. Altogether, the dataset contains 4068 digitized pages. It contains approximately 60,000 words written in Arabic by 1017 journalists (around 60 words for every author) for content autonomous investigation and more than 100,000 Arabic words composed by the same 1017 journalists for text-dependent investigation. Also, it contains around 60,000 words for text-dependent investigation in English. In addition to writer identification, the dataset may be valuable for some other research areas including the distinguishing proof of the gender and handedness of a specific author, and additionally his or her age and nationality. An overview of the QUWI dataset is shown in the Figure 4.1[12].

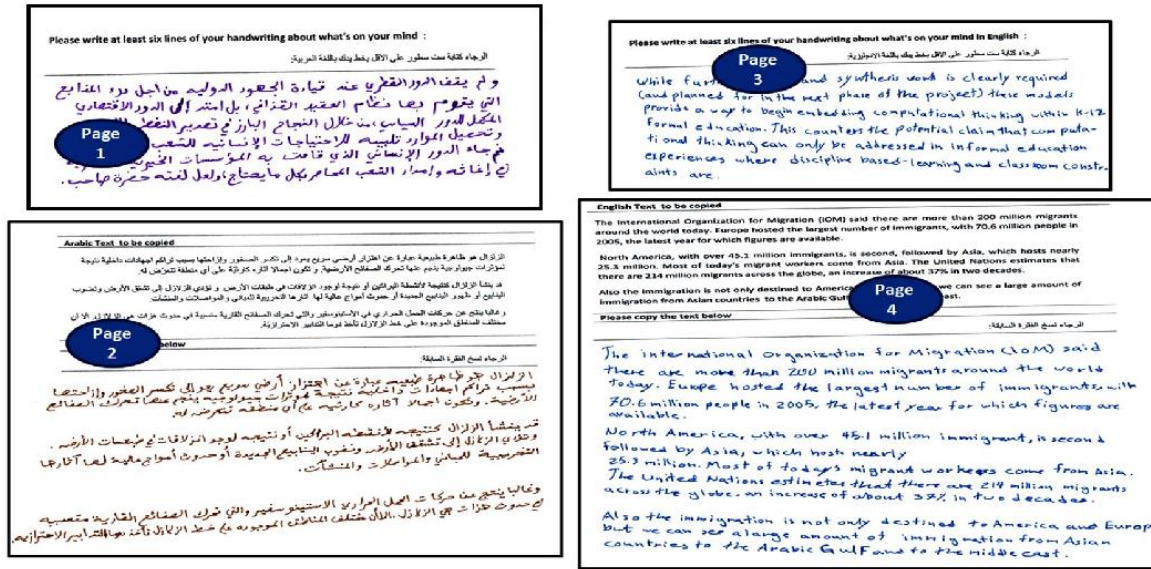


Figure 4.1: Sample Pages of QUWI Dataset

1. First page contains an Arabic handwritten text which varies from one writer to another writer.
2. Second page contains an Arabic H.W text which is the same for all the writers.
3. Third page contains an English handwritten text which varies from one writer to another.
4. Fourth page contains an English H.W text which is the same for all the writers.

We divided the current QUWI dataset into 80% for training and 20% for testing set.

4.2 Proposed Model

We have proposed a system model (Figure: 4.2) for handwritten character recognition and classification system for recognition of the writer's handwriting and predicting the age, gender and handedness. The proposed model goes through the following steps: image acquisition, image preprocessing by removing the duplicate points, edge detection technique for image segmentation, feature extraction and classification.

Image acquisition means getting the image either by scanning documents or by capturing photograph or by directly writing using mouse. Image Preprocessing involves noise removal by removing the duplicate points. The aim of pre-processing is an improvement of the image data

that suppresses unwanted distortions or enhances some image features important for further processing. Image segmentation means to partition the image into multiple segments to make it more meaningful and for easy analyzing. In our model, we have used edge detection for image segmentation. Feature Extraction is said to be a heart of any pattern recognition system. Feature Extraction means to capture relevant characteristic of a target object. It is mainly related to dimension reduction. At first, the features are selected using mRMR (Minimum Redundancy Maximum Relevance) method. The selected features are expected to contain the relevant information from the input data, so that the desired task can be performed by using this reduced representation instead of the complete initial data. The selected features are then extracted using feature extraction method like: -tortuosity, direction, curvature and chain code algorithms. Finally, KNN, SVM & RFC classifiers were used for combining the characterized features and comparing the input with the stored pattern, hence to find best matching class.

4.2.1 Image Acquisition

The training and testing dataset used in this research is an offline dataset called the Qatar University Writer Identification dataset (QUWI). This dataset contains both Arabic and English handwritings and can be used to evaluate the performance of offline writer identification systems. It consists of handwritten documents of 1017 volunteers of different ages, nationalities, genders and education levels. The writers were asked to copy a specific text and to generate a random text, which allows the dataset to be used for both text-dependent and text-independent writer identification tasks.

4.2.2 Dataset Preprocessing

Pre-processing is a common name for operations with images at the lowest level of abstraction - both input and output are intensity images. The aim of pre-processing is an improvement of the image data that suppresses unwanted distortions or enhances some image features important for further processing. We have removed the duplicate points from the dataset in pre-processing step.

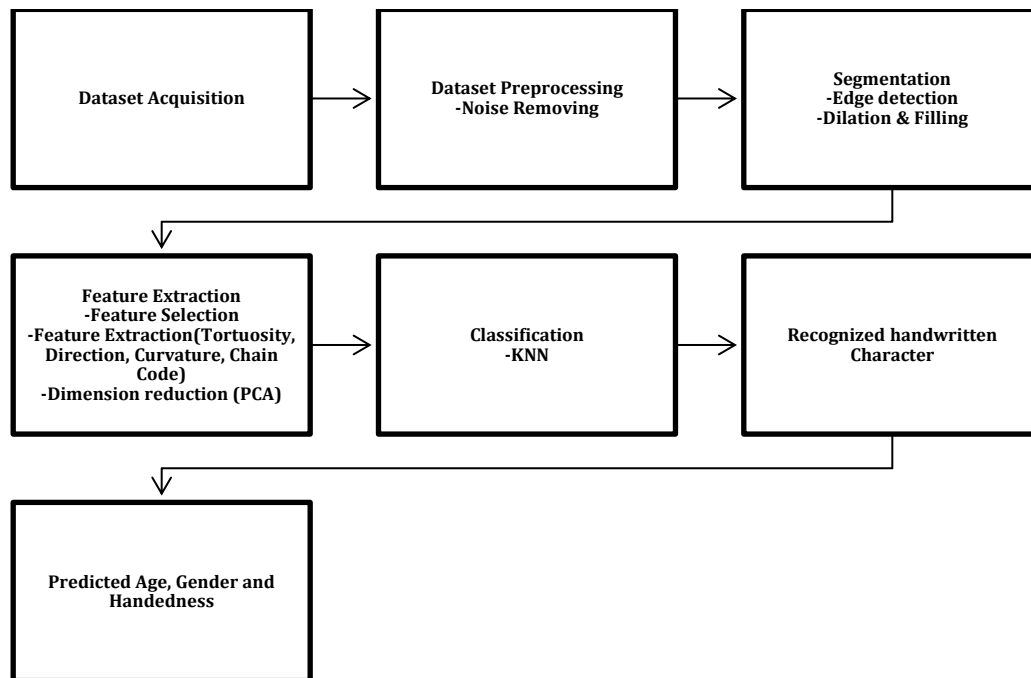


Figure 4.2: Proposed Model

4.2.3 Image Segmentation using Edge Detection Technique

In the image segmentation step, we have used edge detection technique for segmentation of the dataset images. In digital image, the edge is a collection of the pixels whose gray value has a step or roof change, and it also refers to the part where the brightness of the image local area changes significantly. Edge detection mainly identifies object boundaries within images. This includes a variety of mathematical methods that identifies points in a digital image at which the image brightness changes sharply or, has discontinuities. The points at which image brightness changes sharply are typically organized into a set of curved line segments termed edges. Edge detection is an important tool in image processing, machine vision, vision, and in the areas of feature detection and extraction. In this model, we have used fuzzy logic algorithm for edge detection.

Fuzzy logic is a superset of conventional (Boolean) logic that has been extended to handle the concept of partial time values between "completely true" and "completely false" By this definition, fuzzy logic departs from classical two-valued set logic. It uses soft linguistic

system variables and a continuous range of true values in the interval $[0, 1]$, rather than strict binary values. It is basically a multivalued logic that allows intermediate values to be defined between conventional evaluations like yes/no, true/false, etc. Fuzzy logic is also a structured, model-free estimator that approximates a function through linguistic input/output associations. Fuzzy image processing is not a unique theory. Fuzzy image processing is the collection of all approaches that understand, represent and process the images, their segments and features as fuzzy sets.

For edge detection a set of nine pixels, part of a 3×3 or 5×5 window of an image to a set of fuzzy conditions which help to highlight all the edges that are associated with an image. The fuzzy conditions help to test the relative values of pixels which can be present in case of presence on an edge in a gray scale image. This is shown in figure 4.3 :

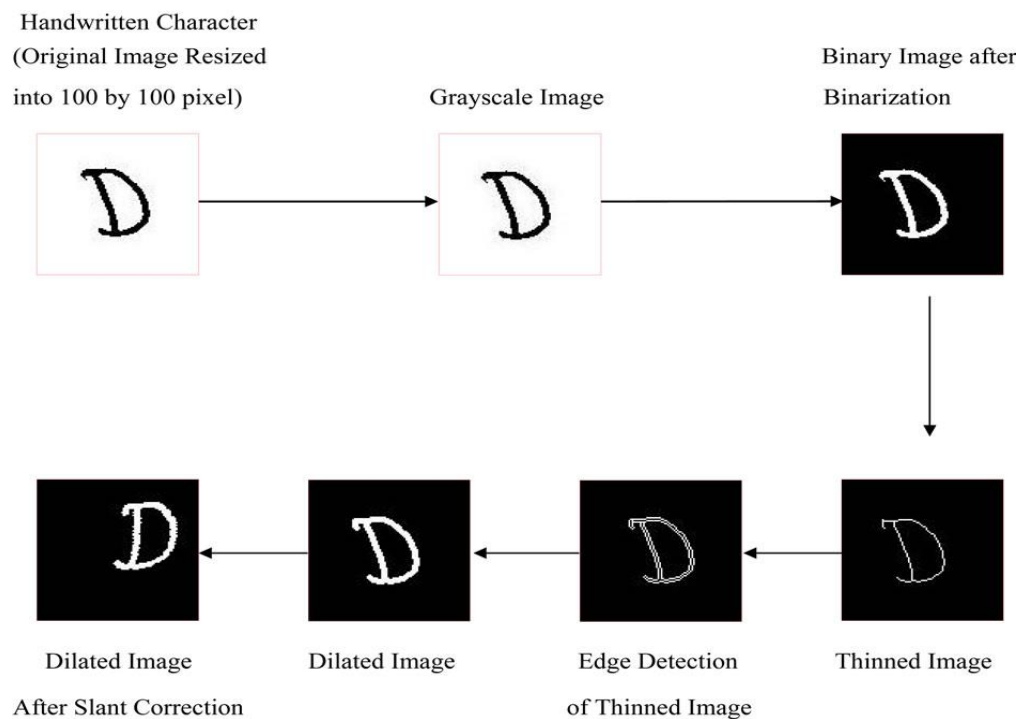


Figure 4.3: Edge Detection Technique

4.2.4 Feature Extraction

We have used mRMR(Minimum Redundancy Maximum Relevance) which is a non-linear feature selection method used for both similarity and importance of characterizing features mRMR use mutual information to measure the importance and the similarity of the features. It uses the mutual information between feature and class as importance and mutual information between the features as similarity [21].

mRMR use the feature selection problem as the below optimization problem [21]:

$$\max_{f_i \in X - S_{m-1}} \left[I(f_i; c) - \frac{1}{m-1} \sum_{f_j \in S_{m-1}} I(f_i; f_j) \right] \dots \dots \dots (1)$$

Here C is class label and this algorithm select one feature which maximizer the equation at every iteration.

Algorithm for Feature Selection with mRMR [21] –

i) Construct a set T to contain the selected features.

Initially T0 = $\emptyset$

ii) Calculate importance of each feature.

iii) Choose most importance feature as first selected feature.

T1=max {imp}

iv) for i=2,....., m :

Choose feature that maximize

$$\max [\alpha \dots imp(f_i) - \beta \cdot \sum \sim (f_i, f_j)] \dots \dots \dots (2)$$

Add feature to t

v) Output T

In the feature extraction step, the characterizing features are extracted from the QUWI handwriting dataset. At first we binarized dataset images using Otsu thresholding algorithm. [13] Then we extracted the characterizing features described below. In the prediction of age, gender and handedness features do not belong to a single value, but a probability distribution function

(PDF) extracted from the dataset images to recognize the identity. The features we consider to predict the age, gender and handedness are following:

- 1) **Tortuosity(F1):** Tortuosity feature distinguish between two writers, fast and slow who produce smooth and twisted handwriting in QUWI dataset, for every pixel P in the text, we took 20 dimensional feature related to tortuosity feature.

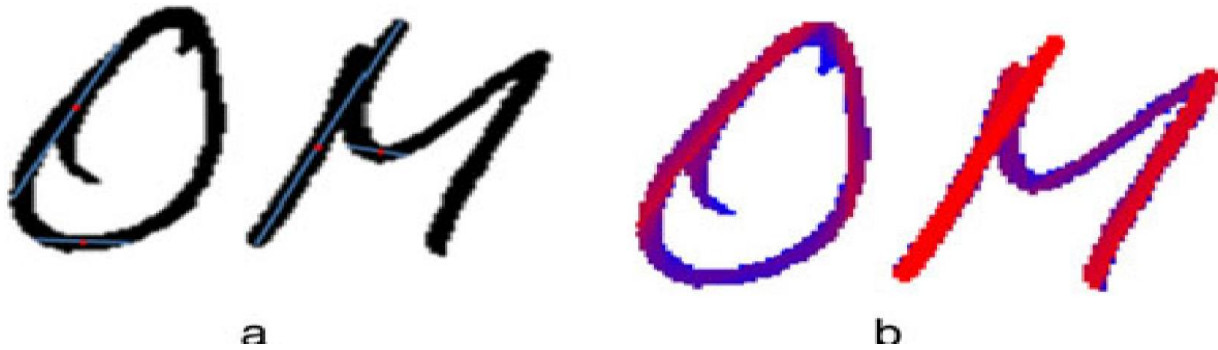


Figure 4.4: Tortuosity Feature

In 20 dimensional features 10 dimensional PDF represents the longest line segment PDF which traverses' P and completed within text and another 10 dimensional PDF represent the directions of the longest line segment. In the figure, longest traversing segment for four different pixels are Length of maximum traversing segment red corresponds to the maximum length, Blue to the minimum one [3].

- 2) **Direction (F2):** Direction feature measures the tangent direction of middle axis of text. Here we used a PDF of 10 dimensions [15]. In the figure, the top figure shows the completed middle axis, then the bottom images show how to find neighborhood pixel for pixels belongs to skeleton which we start from one pixel and travel along preorder directions showed in (a), then a final order pixel is shown in panel (b), after that the final directional at pixel P is computed as the slope of linear regression of 11 neighborhood points as shown in Figure(c)[15].

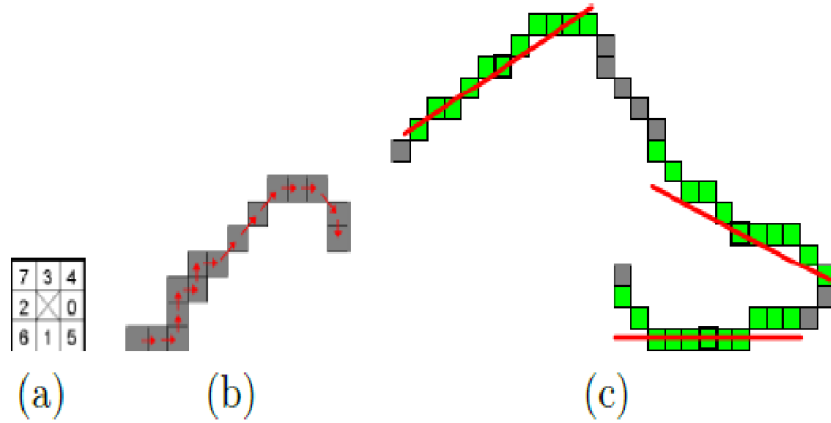


Figure 4.5: Direction feature

- 3) **Chain Code (F3):** Chain codes are produced by browsing the contour of the text and assigning a number to each pixel according to its location with respect to the previous pixel [3]. For every pixel we consider eight possible directions and consider the location. For 1 pixel, we have a 8 dimensional PDF which we used previously, for previous 2 pixel, we have 8^2 dimensional PDF, for 3 pixel, we have a 8^3 dimensional PDF and finally for 4 pixel we have 8^4 dimensional PDF. In the following figure 4.6, we show (a) order followed to produce chain, (b) shape and (c) its corresponding chain code [22].

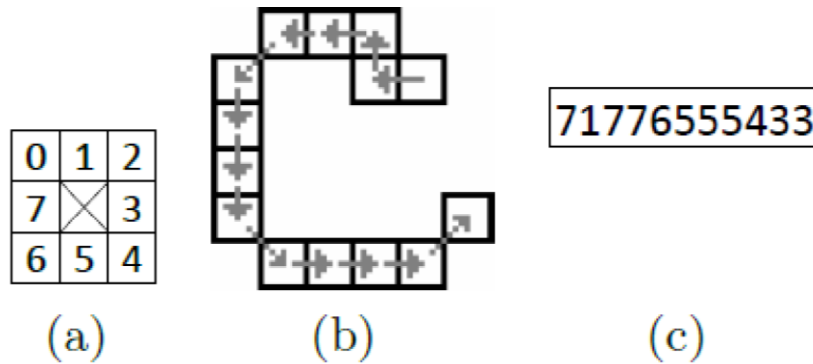


Figure 4.6: Chain Code

- 4) **Curvatures (F4):** Forensic experts consider the curvature as characterizing feature in forensic document examination. PDF of 100 dimensions is used to represent the values of curvature at the contour pixels. For computing each contour pixel, we take a $t=5$ window, and the curvature value is computed as $\frac{(n_1 - n_2)}{n_1 + n_2}$; n_1 is the pixels number belongs the background, while n_2 is the pixel number belongs to foreground shown in the figure 4.7-

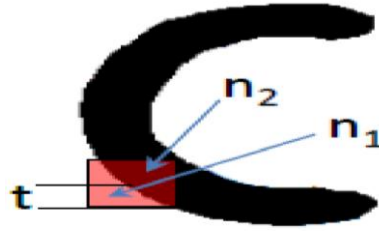


Figure 4.7: Curvatures

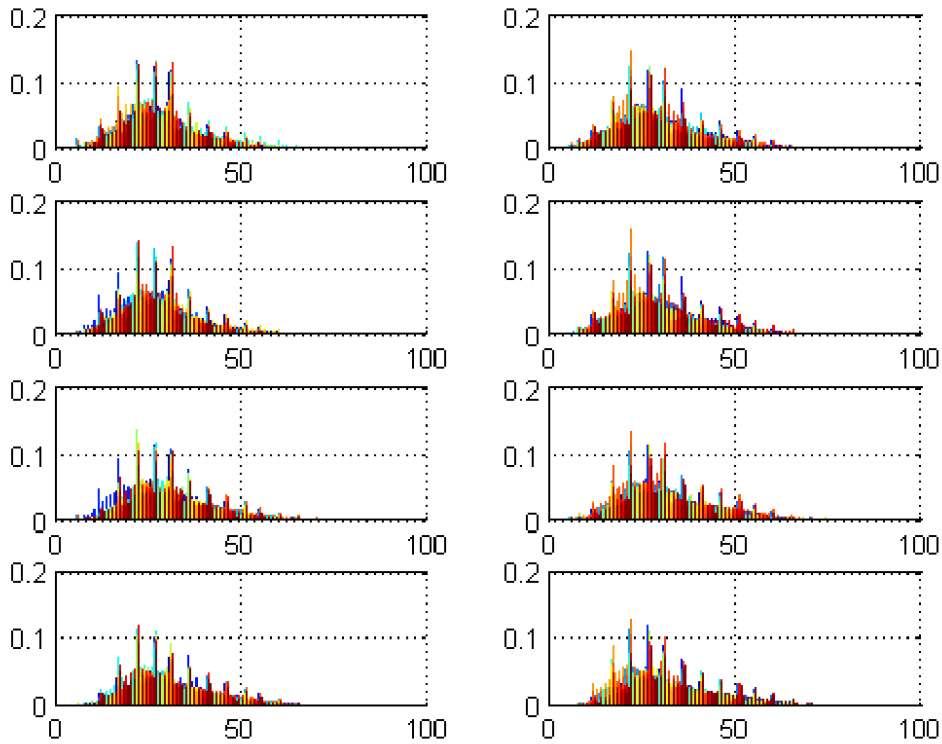


Figure 4.8: Extracted Curvature Value. Left: male; Right: female.

4.2.5 Dimension Reduction using PCA (Principal Component Analysis)

Large datasets are increasingly common and are often difficult to interpret. Principal Component Analysis (PCA) is a technique for reducing the dimensionality of such datasets, increasing interpretability but at the same time minimizing information loss.

Consider a data matrix, X , with column-wise zero experimental mean (the sample mean of each column has been shifted to zero), where each of the n rows represents a different repetition of the experiment, and each of the p columns gives a particular kind of feature (say, the results from a particular sensor).

Mathematically, a set of p -dimensional vectors of weights is used to define the transformation or loadings that map each row vector $x(i)$ of x to a new vector of principal component scores given by $t_{k(i)} = x(i) \cdot w_{(k)}$ for $i = 1, \dots, n$ and $k = 1, \dots, l$... (3) in such a way that the individual variables of t considered over the data set successively inherit the maximum possible variance from x , with each loading vector w constrained to be a unit vector.

First component

The first loading vector $w(1)$ need to satisfy the following equation for maximizing the variance

$$w_{(1)} = \operatorname{argmax}_{||w||=1} \{\sum_i (t1)_{(i)}^2\} = \operatorname{argmax}_{||w||=1} \{\sum_i \dots\dots\dots\} \quad (4)$$

Equivalently, writing this in matrix form gives

$$w_{(1)} = \operatorname{argmax}_{||w||=1} \{|Xw|^2\} = \operatorname{argmax}_{||w||=1} \{W^T X^T X w\} \dots\dots\dots (5)$$

Since $w(1)$ has been defined to be a unit vector, it equivalently also satisfies

$$w_{(1)} = \operatorname{argmax} \left\{ \frac{w^T X^T X w}{w^T w} \right\} \dots\dots\dots (6)$$

The quantity to be maximized can be recognized as a Rayleigh quotient. When w is the corresponding eigenvector, a standard result for a positive semi definite matrix such as $X^T X$ is that the quotient's maximum possible value is the largest eigenvalue of the matrix. With $w(1)$ found, the first principal component of a data vector $x(i)$ can then be given as a score $t1(i) = x(i) \cdot w(1)$ in the transformed co-ordinates, or as the corresponding vector in the original variables, $\{x(i) \cdot w(1)\} w(1)$ [17].

Further components

By subtracting the first $k - 1$ principal component from X , we can find the k th component:

$$\hat{X}_k = X - \sum_{s=1}^{k-1} X w_{(s)} w_{(s)}^T \dots \dots \dots (7)$$

and then finding the loading vector which extracts the maximum variance from this new data matrix

$$w_{(k)} = \underset{|w|=1}{\operatorname{argmax}} \{ |\hat{X}_k w|^2 \} = \underset{|w|=1}{\operatorname{argmax}} \left\{ \frac{w^T \hat{X}_k^T \hat{X}_k w}{w^T w} \right\} \dots \dots \dots (8)$$

With the maximum values for the quantity in brackets given by their corresponding eigenvalues, It turns out that this gives the remaining eigenvectors of $\mathbf{X}^T \mathbf{X}$. Thus, the loading vectors are eigenvectors of $\mathbf{X}^T \mathbf{X}$.

The k th principal component of a data vector $\mathbf{x}_{(i)}$ can therefore be given as a score $t_{k(i)} = \mathbf{x}_{(i)} \cdot \mathbf{w}_{(k)}$ in the transformed coordinates, or as the corresponding vector in the space of the original variables, $\{\mathbf{x}_{(i)} \cdot \mathbf{w}_{(k)}\} \mathbf{w}_{(k)}$, where $\mathbf{w}_{(k)}$ is the k th eigenvector of $\mathbf{X}^T \mathbf{X}$ [17].

The full principal components decomposition of $\mathbf{X}$ can therefore be given as

$$\mathbf{T} = \mathbf{XW}$$

where $\mathbf{W}$ is a p -by- p matrix whose columns are the eigenvectors of $\mathbf{X}^T \mathbf{X}$. The transpose of $\mathbf{W}$ is sometimes called the whitening or sphering transformation.

Covariance

$\mathbf{X}^T \mathbf{X}$ itself can be recognized as proportional to the empirical sample covariance matrix of the dataset $\mathbf{X}$. The sample covariance Q between two of the different principal components over the dataset is given by:

$$\begin{aligned} Q(PC_{(j)}, PC_{(k)}) &\propto (X w_{(k)})^T X w_{(j)} \\ &= w_{(j)}^T X^T X w_{(k)} \\ &= w_{(j)}^T \lambda_{(k)} w_{(k)} \\ &= \lambda_{(k)} w_{(j)}^T w_{(k)} \dots \dots \dots (9) \end{aligned}$$

Where the eigenvalue property of $\mathbf{w}_{(k)}$ has been used to move from line 2 to line 3. However, eigenvectors $\mathbf{w}_{(j)}$ and $\mathbf{w}_{(k)}$ corresponding to eigenvalues of a symmetric matrix are orthogonal (if the eigenvalues are different) or can be orthogonalized (if the vectors happen to share an equal repeated value). As there is no sample covariance between different principal components over the dataset, therefore the product in the final line is zero [17].

Another way to characterize the principal components transformation is therefore as the transformation to coordinates which diagonalize the experimental sample covariance matrix.

In matrix form, the empirical covariance matrix for the original variables can be written [18]

$$Q \propto X^T X = W \Lambda W^T \dots\dots\dots(10)$$

The empirical covariance matrix between the principal components becomes

$$W^T Q W \propto W^T W \Lambda W^T W = \Lambda \dots\dots\dots(11)$$

Where Λ is the diagonal matrix of eigenvalues $\lambda_{(k)}$ of $\mathbf{X}^T \mathbf{X}$

($\lambda_{(k)}$ being equal to the sum of the squares over the dataset associated with each component k : $\lambda_{(k)} = \sum_i t_{k(i)}^2 = \sum_i (\mathbf{x}_{(i)} \cdot \mathbf{w}_{(k)})^2$) [19].

Dimensionality reduction

The transformation $\mathbf{T} = \mathbf{X} \mathbf{W}$ maps a data vector $\mathbf{x}_{(i)}$ from an original space of p variables to a new space of p variables which are uncorrelated over the dataset. However, not all the principal components need to be kept. Keeping only the first L principal components, produced by using only the first L loading vectors, gives the truncated transformation

$$T_L = X W_{L \times p} \dots\dots\dots(12)$$

Where the matrix $\mathbf{T}_L$ now has n rows but only L columns. In other words, PCA learns a linear transformation $t = W^T x, x \in R^p, t \in R^L$, where the columns of $p \times L$ matrix W form an orthogonal basis for the L features (the components of representation t) that are decor related. By construction, of all the transformed data matrices with only L columns, this score matrix maximizes the variance in the original data that has been preserved, while minimizing the total squared reconstruction error [20]

$$\|T W^T - T_L W_L^T\|_F^2 \dots\dots\dots(13)$$

Such dimensionality reduction can be an exceptionally valuable step for imagining and handling high-dimensional datasets, while as yet holding however much of the variance in the dataset as could be expected. For instance, choosing $L = 2$ and keeping just the initial two central parts finds the two-dimensional plane through the high-dimensional dataset in which the information is most spread out, or if the information contains bunches these too might be most spread out, and hence most unmistakable to be plotted out in a two-dimensional outline; though if two bearings through the information (or two of the first factors) are picked indiscriminately, the groups might be considerably less spread separated from each other, and may in reality be significantly more prone to generously overlay each other, making them indistinguishable.

Dimensionality reduction may likewise be fitting when the factors in a dataset are boisterous. On the off chance that every segment of the dataset contains free indistinguishably appropriated Gaussian commotion, at that point the sections of T will likewise contain correspondingly indistinguishably disseminated Gaussian noise (such a dispersion is invariant under the impacts of the grid W , which can be thought of as a high-dimensional revolution of the facilitate tomahawks). Be that as it may, with a greater amount of the aggregate fluctuation moved in the initial couple of chief segments contrasted with a similar commotion change, the proportionate impact of the noise is less—the initial couple of segments accomplish a higher flag-to-clamor proportion. PCA in this manner can have the impact of concentrating a significant part of the flag into the initial couple of chief segments, while the later principal components may be dominated by noise, and so disposed of without great loss.

4.3 Classification:

To show working accuracy of our system model and combining the extracted features we have used three classifiers as follows:

- K-Nearest Neighbor(KNN)
- Random Forest Classifier(RFC)
- Support Vector Machine(SVM)

Feature extraction methods are described above to take decision category of our handwriting writer. In the classification step every feature vectors are used as an individual input

for our classifier. We combined Tortuosity, Direction and Chain Code. Curvatures feature extraction methods using KNN, RFC and SVM. Description of these classifiers is given below:

1) K-Nearest Neighbor(KNN)

The KNN algorithm is a robust and versatile classifier that is often used as a benchmark for more complex classifiers such as Artificial Neural Networks (ANN) and Support Vector Machines (SVM). Despite its simplicity, KNN can outperform more powerful classifiers and is used in a variety of applications such as economic forecasting, data compression and genetics. The KNN classifier is a non-parametric and instance-based learning algorithm. The advantages of using KNN are that it is robust to noisy training data and effective if the data is large.

To work well, this algorithm requires a training dataset which is a set of well labeled data points. This algorithm takes as input a new data point and makes the classification for this by calculating the Euclidean or Hamming distance between the new data point and the labeled data point [5]. The Euclidean distance is calculated using the following formula:

$$d(p, q) = d(p, q) = \sqrt{(q_1 - p_1)^2 + (q_2 - p_2)^2 + \dots + (q_n - p_n)^2} \\ = \sqrt{\sum_{i=1}^n (q_i - p_i)^2} \dots \dots \dots (14)$$

2) Random Forest Classifier(RFC)

Random Forest algorithm is a supervised classification algorithm that implies that there is a direct relationship between the number of trees in the forest and the results it can get: the larger the number of trees, the more accurate the result. It can be used for both classification and regression tasks. For Random Forest algorithm, if there are enough trees in the forest, the classifier won't over fit the model, thus avoiding the over fitting problem. The third advantage is the classifier of Random Forest can handle missing values, and the last advantage is that the Random Forest classifier can be modeled for categorical values. To make predictions, once the training is done, the average of predictions from all individual regression trees is taken using the following formula:

$$\hat{f} = \frac{1}{B} \sum_{b=1}^B \hat{f}_b(X') \dots \dots \dots (15)$$

3) Support Vector Machine(SVM)

“Support Vector Machine” (SVM) is a supervised machine learning algorithm which can be used for both classification and regression challenges. In this algorithm, each data item is plotted as a point in n-dimensional space (where n is number of features) with the value of each feature being the value of a particular coordinate. Then, classification is performed by finding the hyper-plane that differentiates the two classes very well. There are four main advantages of this algorithm: firstly it has a regularization parameter, which makes the user think about avoiding over-fitting. Secondly it uses the kernel trick, so you can build in expert knowledge about the problem via engineering the kernel. Thirdly an SVM is defined by a convex optimization problem (no local minima) for which there are efficient methods (e.g. SMO). Lastly, it is an approximation to a bound on the test error rate, and there is a substantial body of theory behind it which suggests it should be a good idea.

Chapter 5

Experimental Results

We conducted experiment for every individual feature which were described on previous chapter and we classified them using RFC, KNN, SVM algorithms. After classification we analyzed and discuss the result. To perform the classification 80% of dataset has been used for training and 20% of dataset for testing.

5.1 Classification accuracy for Age

For age classification 5 ranges were defined for a good accuracy (1950 to 1959, 1960 to 1969, 1970 to 1979, 1980 to 1989 and 1990 to 1999). Here, tortuosity, direction, chain code and curvature give average accuracy of 84.02%, 86.79%, 86.53% and 84.91% accordingly. Combining tortuosity with direction gives an average of 87.14% accuracy, chain code and curvature combined gives average 85.87% accuracy, tortuosity and curvature jointly gives

Table 5.1: Classification accuracy for age

Feature	Arabic (%)			English (%)			Both (%)		
	KNN	SVM	RFC	KNN	SVM	RFC	KNN	SVM	RFC
F1	83.7	86.9	82.6	82.6	83.7	81.2	85.2	84.1	86.2
F2	86.7	89.8	85.7	85.6	87.8	84.3	87.8	86.2	87.23
F3	86.5	84.2	86.7	84.5	88.8	86.8	85.7	87.3	88.32
F4	82.3	83.1	81.5	85.5	87.6	87.2	84.6	85.2	87.23
F1+F2	87.4	86.7	88.3	88.2	91.5	85.5	83	90.3	83.43
F3+F4	88.8	87.3	88.2	86.5	86.4	84.3	83.3	82.5	85.55
F1+F4	86.8	82.6	84.1	88.1	91.3	88.4	89.8	89.3	83.43
F2+F4	89.7	92.9	87.2	91.1	95.4	91.5	92.4	91.4	89.46
F1+F2+F3+F4	89.2	94.6	91.5	94.7	97.9	82.3	93.3	92.4	91.98

average 87.09% accuracy and direction and curvature combination gives average 91.22% accuracy. Finally, combining all the features, we get an accuracy of average 91.97%.

5.2 Classification accuracy for Gender

For gender classification, a random classifier would predict approximately 50%, as this is a two-class or binary classification. Here, tortuosity gives an accuracy of average 77.02%, direction gives average 76.96%, chain code gives an average of 73.17% and curvature gives average 75.12%. Combining tortuosity with direction gives an average of 76.04% accuracy, chain code and curvature combinedly gives average 73.76% accuracy, tortuosity and curvature jointly gives average 74.59% accuracy and direction and curvature combination gives average 76.49% accuracy. Finally, combining all the features, we get an accuracy of average of almost 80%. For gender detection form Arabic character, KNN, SVM and RFC shows average 77.43%, 76.97% and 75.22% accuracy. Whereas, for gender detection form English character, KNN, SVM and RFC shows average 75.02%, 78.69% and 73.79% accuracy. And for both Arabic and English character, the average accuracy rate is 74.47%, 76.38% and 75.18% for KNN, SVM and RFC accordingly.

Table 5.2: Classification accuracy for gender

Feature	Arabic (%)			English (%)			Both (%)		
	KNN	SVM	RFC	KNN	SVM	RFC	KNN	SVM	RFC
F1	78.7	77.8	76.8	76.5	75.2	74.3	77.8	76.9	79.2
F2	77.8	76.7	75.2	77.2	81.3	71.3	76.5	80.4	76.2
F3	76.23	67.2	69.3	78.2	80.2	68.23	69.3	78.7	71.2
F4	76.3	76.2	73.2	71.2	78.3	73.2	77.6	76.2	73.9
F1+F2	77.5	75.3	75.2	72.3	77.2	74.2	77.2	77.2	78.3
F3+F4	75.3	76.2	76.7	71.2	75.3	75.2	69.5	69.2	75.2
F1+F4	77.24	79.81	72.2	73.4	78.5	73.2	71.45	72.33	73.2
F2+F4	78.3	80.12	78.2	75.6	79.7	74.2	73.44	77.23	71.6
F1+F2+F3+F4	79.5	83.4	80.2	79.56	82.5	80.32	77.45	79.26	77.8

5.3 Classification accuracy for Handedness

Here, tortuosity gives an accuracy of average 71.98%, direction gives average 72.65%, chain code gives an average of 72.83% and curvature gives average 75.18%. Combining tortuosity with direction gives an average of 77% accuracy, chain code and curvature combinedly gives average 76.1% accuracy, tortuosity and curvature jointly gives average 74.05% accuracy and direction and curvature combination gives average 77.68% accuracy. Finally, combining all the features, we get an accuracy of average of 78.53%.

For Left Handedness detection, KNN, SVM and RFC show average 65.51%, 64.39% and 64.06% accuracy. Whereas, for Right Handedness detection, KNN, SVM and RFC shows average 84.46%, 86% and 84.27% accuracy.

Table 5.3: Classification accuracy for handedness

Feature	Left Handed (%)			Right Handed (%)		
	KNN	SVM	RFC	KNN	SVM	RFC
F1	61.2	60.2	64.2	80.6	82.5	83.2
F2	63.7	59.3	62.3	83.7	83.6	83.3
F3	67.8	58.2	61.4	82.6	82.6	84.4
F4	64.4	59.2	64.3	88.8	88.8	85.6
F1+F2	66.2	67.2	66.2	85.6	89.2	87.6
F3+F4	64.3	68.2	64.3	86.6	86.6	86.6
F1+F4	67.2	69.8	62.2	93.4	78.5	73.2
F2+F4	68.3	70.1	68.2	85.6	89.7	84.2
F1+F2+F3+F4	66.5	67.3	63.4	91.2	92.5	90.3

5.4 Classification accuracy for Handwritten Character recognition

To recognize the characters, we combined all the features and used three classifiers KNN, SVM and RFC. After the training of the dataset, we were able to recognize all the 26 alphabets successfully. Averaging the accuracy of KNN, SVM and RFC, we get 98.2% for A, 80.22% for B, 99.77% for C, 98.81% for D, 93.57% for E, 96.34% for F, 91.63% for G, 76.91% for H,

Table 5.4: Classification accuracy for recognizing handwritten characters (English)

Feature	Character	KNN	SVM	RFC
F1+F2+F3+F4	A	99.6	96.5	98.5
F1+F2+F3+F4	B	79.86	78.5	82.3
F1+F2+F3+F4	C	100	100	99.3
F1+F2+F3+F4	D	100	97.2	99.23
F1+F2+F3+F4	E	89.2	95.2	96.3
F1+F2+F3+F4	F	96.5	97.3	95.23
F1+F2+F3+F4	G	92.3	93.39	89.2
F1+F2+F3+F4	H	76.3	73.23	81.2
F1+F2+F3+F4	I	79.56	82.35	81.3
F1+F2+F3+F4	J	97.8	96.2	99.3
F1+F2+F3+F4	K	85.7	87.2	83.2
F1+F2+F3+F4	L	92.5	91.3	94.5
F1+F2+F3+F4	M	85.3	88.5	83.4
F1+F2+F3+F4	N	86.4	87.01	91.3
F1+F2+F3+F4	O	80.1	76.23	79.3
F1+F2+F3+F4	P	91.3	95.6	94.4
F1+F2+F3+F4	Q	66.3	71.5	68.3
F1+F2+F3+F4	R	73.4	71.5	78.8
F1+F2+F3+F4	S	95.6	90.3	97.7
F1+F2+F3+F4	T	91.3	93.4	96.5
F1+F2+F3+F4	U	89.4	83.5	88.5
F1+F2+F3+F4	V	98.04	96.7	99.5
F1+F2+F3+F4	W	91.5	90.30	94.5
F1+F2+F3+F4	X	85.07	86.70	89.56
F1+F2+F3+F4	Y	91.5	96.43	93.2
F1+F2+F3+F4	z	95.3	92.4	91.5
	Over all %	89.2381481	88.7861538	90.2315385

81.07% for I, 97.77% for J, 85.37% for K, 92.77% for L, 85.73% for M, 88.27% for N, 78.54% for O, 93.57% for P, 68.7% for Q, 74.57% for R, 94.53% for S, 93.73% for T, 87.13% for U,

98.08% for V, 92.1% for W, 87.11% for X, 93.71% for Y and 93.07% for Z. For recognizing all the character, RFC showed the best average accuracy of 90.23%, followed by 89.24% for KNN and 88.77% for SVM.

5.4 Performance analysis

We have analyzed the average performance for every feature (Figure 5.2) we examined for gender, age & handedness. For age, tortuosity gave 84.02%, direction gave 86.79%, chain code gave 86.53% & curvatures gave 84.91% accuracy. For gender tortuosity gave 77.02%, direction gave 76.95%, chain code gave 73.13% & curvatures gave 75.12% accuracy. For handedness, tortuosity gave 71.98%, direction gave 72.65%, chain code gave 72.83 % & curvatures gave 75.18% accuracy.

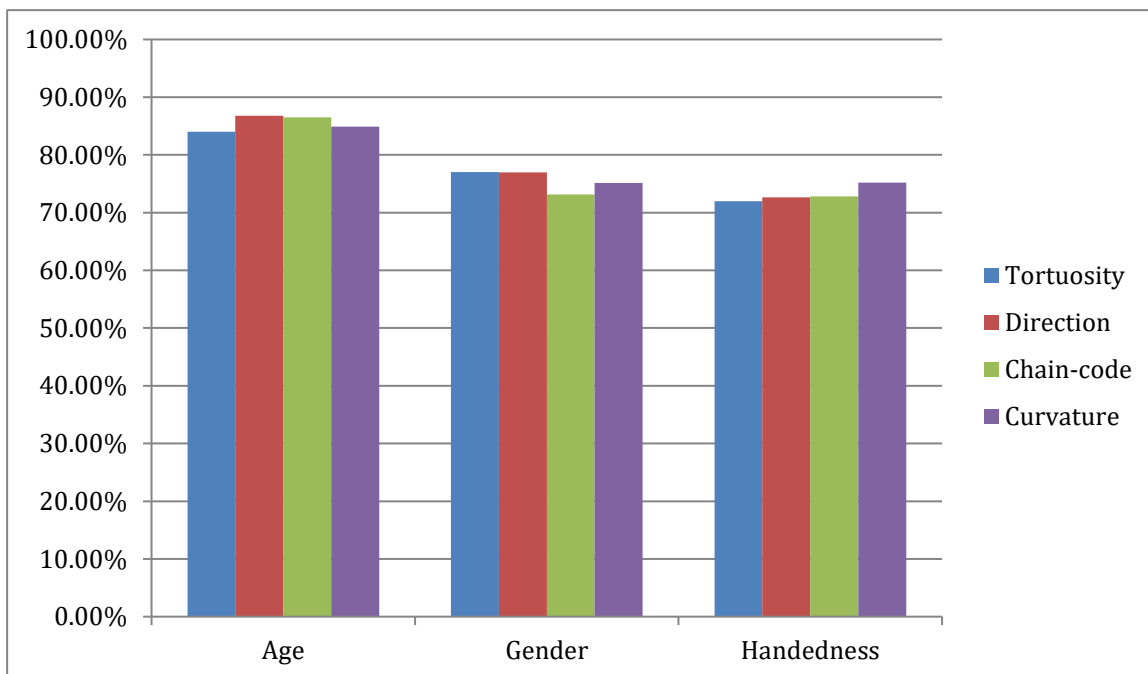


Figure 5.2: Performance analysis of every feature for Predicting age, gender & handedness

For recognizing all the character, RFC showed the best average accuracy of 90.23%, followed by 89.24% for KNN and 88.77% for SVM. Which means RFC can recognize the alphabets more accurately than other classifiers.

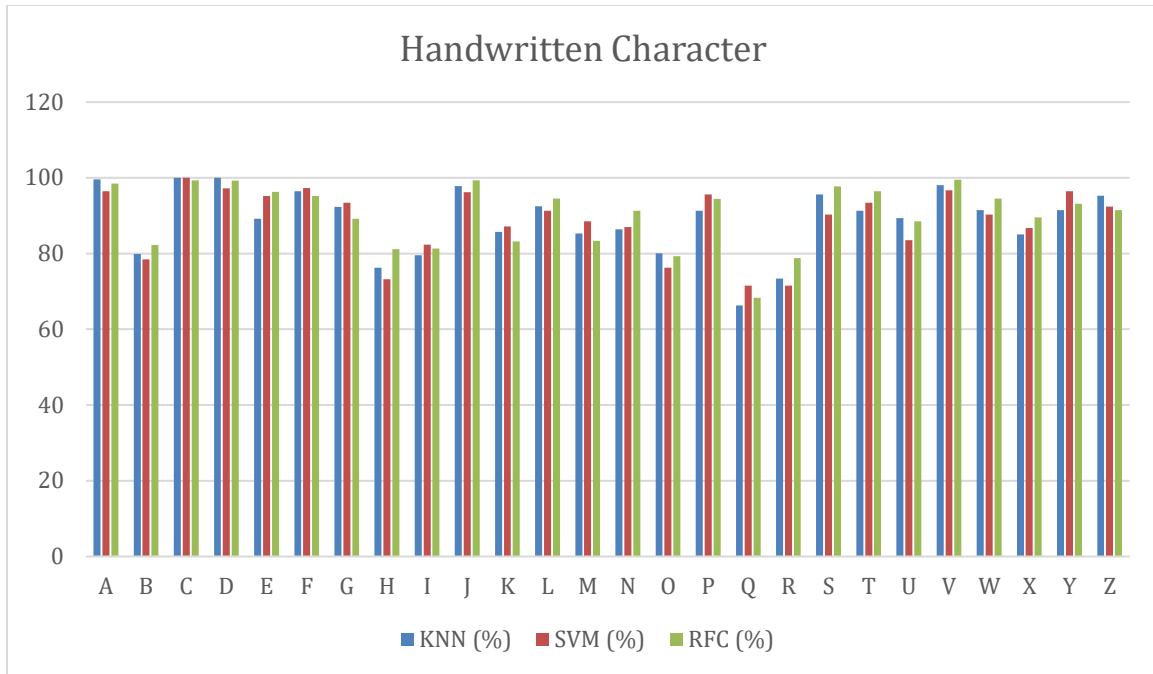


Figure 5.3: Performance analysis of every feature for handwritten character recognition.

5.5 Comparative Study

S.Al-ma'adeed and A.Hassaine in their paper [3] combined the features using random forests and kernel discriminant analysis. Classification accuracy are reported on the QUWI dataset, reaching 74.05% for gender prediction and 55.76% for age range prediction. N. Bouadjenek, H. Nemmour and Y. Chibani [24] conducted a research for predicting age, gender and handedness using gradient features. For the classification step, authors employed random forest and kernel discriminant analysis classifiers. Obtained accuracy are about 73.59%, 60.62% and 50% for gender, age range and nationality classification.

In our paper, we used KNN, SVM, RFC algorithms for finding the accuracy and we used QUWI dataset. Our average accuracy is 88.28% for age range prediction, 75.90% for gender prediction and 75.11% for handedness prediction for each writer. So, comparing the above mentioned paper our accuracy rate is better than their accuracy.

Chapter 6

Conclusion and Future Work

In this paper, we have tried to come up with a solution to perform two tasks simultaneously and efficiently: recognize handwritings from various input sources/devices and make a prediction of the author's age, gender and handedness. Artificial intelligence offers many benefits in pattern recognition and classification within the sense of emulating accommodative human intelligence to a little extent Handwriting Recognition is the first step to the vast field of Artificial Intelligence and Image Processing. With the advancement of technology, machines can now identify images and understand handwriting of different people which humans themselves have failed to comprehend from the image. The extracted features are used to train an artificial neural network which classifies the demographics of the author of a query. Example to accurately classify a cursive handwriting or different styles of handwriting is tough, even for an expert. Thus, keeping all these factors and the complexity of the situation in mind we have thus decided to come up with this solution. This model shall be a new step towards digitalization and is believed to have an immense contribution in this field. In future, we would like to extend our model to read Bengali handwritings and contribute to the Bengali community where not much work has been done.

Chapter 7

References

1. M. Liwicki, A. Schlappbach, P. Loretan, and H. Bunke. Automatic Detection of Gender and Handedness from On-Line Handwriting, *EURASIP Journal on Image and Video Processing*, December 2014, vol. 10, pp. 1-10, 2014
2. A. Dutta and A. Dutta, "Handwritten digit recognition using deep learning, "International Journal of Advanced Research in Computer Engineering & Technology (IJARCET)," vol. 6, no. 7, July 2017.
3. S. Al. Maadeed and A. Hassaine. (2014). Automatic prediction of age, gender, and nationality in offline handwriting. *EURASIP Journal on Image and Video Processing*. [Online]. Available: <https://doi.org/10.1186/1687-5281-2014-10>
4. Abdul Hamid, Norhidayu & Nur Amir Sjarif, Nilam. (2017). Handwritten Recognition Using SVM, KNN and Neural Network.
5. Ángel Morera, Ángel Sánchez, José Francisco Vélez, and Ana Belén Moreno, "Gender and Handedness Prediction from Offline Handwriting Using Convolutional Neural Networks," *Complexity*, vol. 2018, Article ID 3891624, 14 pages, 2018. <https://doi.org/10.1155/2018/3891624>.
6. G. Jain and J. Ko, "Handwritten digits recognition", Project report University of Toronto, 2008.
7. Java OCR, 5 June 2010. Retrieved 5 June 2010. [Online]. Available: <http://roncemer.com/software-development/java-ocr/>
8. Chakravarthy, A. S. N., Penmetsa V. Krishna Raja and P. S. Avadhani. "Handwritten Text Image Authentication using Back Propagation." *CoRR abs/1110.1488* (2011): n. pag. (link <https://arxiv.org/ftp/arxiv/papers/1110/1110.1488.pdf>)
9. Huang, B.; Zhang, Y. and Kechadi, M.; Preprocessing Techniques for Online Handwriting Recognition. *Intelligent Text Categorization and Clustering*, Springer Berlin Heidelberg, 2009, Vol. 164, "Studies in Computational Intelligence" pp. 25–45.
10. Holzinger, A., Stocker, C. Peischl, B. and Simonc, K. M., On Using Entropy for Enhancing Handwriting Preprocessing, *Entropy* 2012, 14, pp. 2324-2350.

11. Y. LeCun, C. Cortes. The MNIST database of handwritten digits. [Online]. Available: <http://yann.lecun.com/exdb/mnist>
12. QUWI: An Arabic and English Handwriting Dataset for Offline Writer Identification
Somaya Al Maadeed, Wael Ayouby, Abdel aali Hassaine, Jihad Mohamad Aljaam
13. N Otsu, A threshold selection method from gray-level histograms. IEEE Trans. Syst. Man Cybern. 9(1), 62–66 (1979)
14. N. Bouadjenek, H. Nemmour, Y. Chibani, "Local descriptors to improve off-line handwriting-based gender prediction," In Proceedings of the 6th International Conference of Soft Computing and Pattern Recognition (SoCPaR), 2014, pp.43 - 47, 11-14 Aug. 2014.
15. J Matas, Z Shao, J Kittler, Estimation of curvature and tangent direction by median filtered differencing, in The 8th International Conference on Image Analysis and Processing. Lecture notes in computer science. vol 974 (Springer-Verlag, Berlin, 1995), pp. 83–88. 13–15 September.
16. A. Hassaïne, S. Al-Maadeed and A. Bouridane. "A set of geometrical features for writer identification." Neural Information Processing. Springer Berlin Heidelberg, 2012.
17. Available[Online]: https://en.wikipedia.org/wiki/Principal_component_analysis
18. Bengio, Y.; et al. (2013). "Representation Learning: A Review and New Perspectives" (PDF). Pattern Analysis and Machine Intelligence. 35 (8): 1798–1828. arXiv:1206.5538 Freely accessible. doi:10.1109/TPAMI.2013.50
19. A. A. Miranda, Y. A. Le Borgne, and G. Bontempi. New Routes from Minimal Approximation Error to Principal Components, Volume 27, Number 3 / June, 2008, Neural Processing Letters, Springer
20. Hotelling, H. (1933). Analysis of a complex of statistical variables into principal components. Journal of Educational Psychology, 24, 417–441, and 498–520.
21. Shirzad, M. and Keyvanpour, M. (n.d.). A feature selection method based on minimum redundancy maximum relevance for learning to rank.
22. Xie, Q. and Xu, Q. Gender Prediction from Handwriting.
23. K. Reddy, D. Rao, U. Ankaiah and K. Rajesh, "Handwritten Character and Digit Recognition Using Artificial Neural Networks", International Journal of Advanced

Research in Computer Science and Software Engineering, vol. 3, no. 4, pp. 140-143, 2013.

24. N. Bouadjenek, H. Nemmour and Y. Chibani, "Age, gender and handedness prediction from handwriting using gradient features," *2015 13th International Conference on Document Analysis and Recognition (ICDAR)*, Tunis, 2015, pp. 1116-1120. doi: 10.1109/ICDAR.2015.7333934