

A robust script independent handwriting system for gender identification

Shivakumara Palaiahnakote^{a,*}, Maryam Asadzadeh Kaljahi^b, Swati Kanchan^c, Umapada Pal^c, Daniel Lopresti^d, Tong Lu^e

^a School of Science, Engineering and Environment, University of Salford, Manchester, UK

^b promiseQ, Berlin, Germany

^c Computer Vision and Pattern Recognition Unit, Indian Statistical Institute, Kolkata, India

^d Computer Science & Engineering, Lehigh University, Bethlehem, PA, USA

^e National Key Lab for Novel Software Technology, Nanjing University, Nanjing, China

ARTICLE INFO

Keywords:

Gradient image
Word segmentation
Gabor orientation
Gabor scaling
Correlation coefficient
Mahalanobis distance
Gender identification

ABSTRACT

Gender identification at the word level in a multi-script environment is challenging due to variations posed by free-style handwriting of individuals and geographical differences in writing styles. This paper presents a new approach, Multi-Orientation-Scale Gabor Response Fusion (MOSGF), for gender identification at the word level using handwritten text. Our method has two steps: (i) word segmentation from unconstrained lines and (ii) gender identification at the word level. In the first step, the method explores the number of zero crossing points and gradient information for word segmentation from handwritten text lines. In the second step, employs Gabor responses at different orientations and scales to detect fine details in female and male handwriting. For each Gabor response, the proposed model estimates the correlation between average templates of all Gabor responses and the individual Gabor response to extract global consistency in writing. To strengthen correlation features, the proposed method uses the Mahalanobis distance measure, which extracts local similarity. Further, the proposed approach fuses correlation coefficient and distance-based features in a novel way. The fused features are then fed to a Neural Network (NN) for gender identification. Experiments on our dataset, which comprises Roman (English), Chinese, Farsi (Persian), Arabic, and Indian scripts, and a benchmark dataset, namely, IAM which includes English text, KHATT which includes Arabic, and QUWI which includes both English and Arabic, show that the proposed system outperforms the existing methods in terms of word segmentation and gender identification.

1. Introduction

Gender identification has received special attention from several applications like forensic analysis where biometric features can be used to help to identify individuals of interest (Manyala et al., 2018). In this area, handwriting analysis plays a pivotal role because writing is often more prevalent and easier to capture, and process compared to other biometric modalities. Handwriting-based gender identification may help focus an investigation on a particular group of individuals (Bi et al., 2018). Additionally, it can also be useful for historical document authentication (Topaloglu and Ekmekci, 2017). Turning again to the broader question, it is noted in the literature that there are a number of approaches to gender identification such as (i) Graphology-based (Topaloglu and Ekmekci, 2017), which uses character shapes and strokes, (ii) Non-biometric features, such as text, speech, clothing and

hair style, (iii) Soft biometric features such as height, weight, color of the eyes color, silhouette, age, gender, race, moles, tattoos, birthmarks, scars, etc, and (iv) Biometric features, such as face, iris, ear, fingerprint, voice, gait, gesture, lip motion and writing style, etc. These features are not always robust in real-world environments. Graphology-based approaches perform well for particular applications because the hypothesis and conditions are framed experimentally but not scientifically (Navya et al., 2018a). In the same way, non-biometric, soft biometric and biometric based approaches are sensitive to issues that arise in open environments.

As a result, to assist forensic investigations, handwriting-based gender identification can make a useful contribution. Such approaches are generally more stable and reliable compared to the above-mentioned approaches. Several such methods have been described in the literature (Boudajenek et al., 2016). However, the scope of these past approaches

* Corresponding author.

E-mail addresses: s.palaiahnakote@salford.ac.uk (S. Palaiahnakote), maryam@promiseq.com (M.A. Kaljahi), swati.kanchan@razorpay.com (S. Kanchan), umapada@isical.ac.in (U. Pal), lopresti@cse.lehigh.edu (D. Lopresti), lutong@nju.edu.cn (T. Lu).

<https://doi.org/10.1016/j.eswa.2024.123576>

Received 19 May 2021; Received in revised form 19 February 2024; Accepted 25 February 2024

Available online 29 February 2024

0957-4174/Crown Copyright © 2024 Published by Elsevier Ltd. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

is limited to full pages of handwritten text and full text lines, but not individual words and hence the methods are computationally expensive. Further, free style writing, multi-script, and multi-orientation make the problem of gender identification at the word level challenging and interesting. Thus, this study focuses on a new method for addressing these challenges. Intuitively, some individuals write with more care to be legible, while others are sloppier with their writing, and these differences may have some gender dependence. It can be seen from the sample handwritten word images of different scripts in Fig. 1, where one can notice the differences in terms of geometrical shapes of writing style of female and male writings.

The following are the key contributions of the work reported in this paper.

- Exploring zero crossing points and gradient information for segmenting handwritten words regardless of scripts, orientation, style, and touching.
- Exploiting Gabor responses at different orientations and scales for enhancing fine details in handwritten words images, which provides clues for differentiating female and male handwritten text.
- Fusing correlation coefficients and Mahalanobis distance-based features for studying global and local consistency in writing style at word level for gender identification with the help of neural network classifier.

The rest of the paper is organized as follows. A comprehensive review of handwritten word and text line segmentation, as well as gender identification is presented in Section 2. Section 3 discusses our methods for word segmentation and gender identification in a multi-script environment. Experimental results to validate key steps in our procedures using several different datasets are analyzed in Section 4. Section 5 provides a summary of the work and discusses topics for future research.

2. Related work

Since our aim is to develop a robust system for gender identification using handwritten text at the word level, this section presents a review of the relevant literature in these areas.

2.1. Words/Text lines segmentation

Wshah et al. (2009) developed a method for segmenting words from Arabic text line images using contour and skeleton information. Their approach employs connected component analysis by splitting each text component into smaller components. Then it uses the smaller components to derive lexicons, and then lexicons are used to group the words in text lines. The method is confined to a particular script because the derived lexicons may not work for other scripts. Louloudis et al. (2009) used a conventional approach to segment text lines through binarization and connected component analysis. The distance is estimated between character components to connect the character component as words. It works based on the assumption that the distance between characters is smaller than the distance between the words. This idea works well for printed text lines, but not for handwritten text where one cannot expect

regular spacing between characters and words. Osman (2013) explored contour traversing for segmenting Arabic text lines in handwritten documents. Their approach uses the direction of contours for separating text lines. Since the method involves several heuristics and conditions which are derived based on knowledge of Arabic script, their approach may not be effective for other scripts. Banumathi and Chandra (Banumathi et al., 2016) used line and word spacing for segmenting text lines and words from handwritten Kannada document images. For finding spaces between words and text lines, their approach uses a projection profile analysis. This method performs well for high quality images and a limited number of scripts because the binarization fails for degraded images. Khare et al. (2018) derived weights using gradient values for segmenting handwritten text lines. The weighted gradient values are used for classifying pixels representing ascenders and descenders into one class, and other pixels into another class. This approach works well for text line segmentation, but not word segmentation. Lamsaf et al. (2019) proposed an approach for handwritten text line and word segmentation for Arabic handwriting. Their method performs binarization for the input images using a global thresholding technique. Then it classifies the space between the connected components into two classes, namely, the spaces between words and the spaces between connected components within words. The threshold is fixed based on distance between the connected components when segmenting words. However, this approach may not be robust for handwriting in other scripts as it is based on Arabic script.

It is observed from the above review that most methods follow a conventional paradigm consisting of binarization followed by connected component analysis for text line and word segmentation. The performance of the binarization step depends on the threshold values. In addition, the scope of the methods may be limited to a particular script because it employs knowledge of the script to achieve better results. This is a significant limitation of existing segmentation methods. With this in mind, a new method is proposed for segmenting words which aims to be robust for multiple scripts, orientations, and other degradations that arise in free-style writing. Inspired by the idea presented in (Khare et al., 2018), where weighted gradients are used for segmenting text lines, our work adapts this approach for segmenting words in handwritten text lines.

2.2. Gender identification

Maji et al. (2015) extracted the Euler number from handwritten signatures and use neural networks for gender classification. In order to extract the Euler number, their method binarizes the input images. This is a drawback because the binarization will not work well for degraded document images. The success of the classification depends on the success of the binarization results. Bouadjene et al. (2015) explored gradient features for classification of writer age, gender, and handedness. To improve the performance of the classification, the approach uses HOG, LBP and pixel density features. But features like LBP and HOG are sensitive to pixel values, and hence the method may not be robust across handwritten text produced under different situations. Mirza et al. (2016) extracted texture features using Gabor filters for gender identification. Since the texture of handwriting changes as the script changes,

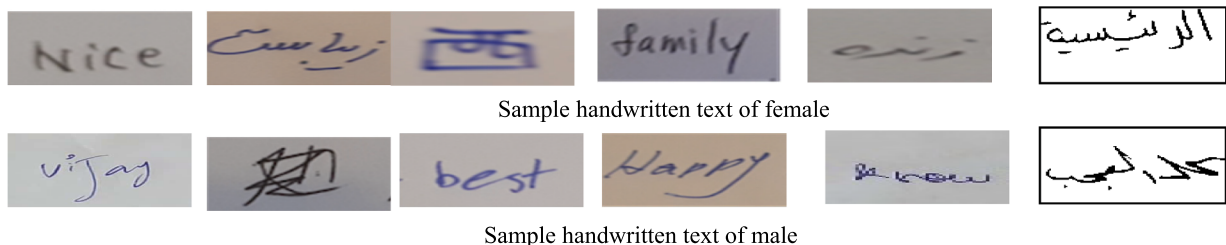


Fig. 1. Clues for separating handwritten text of female and male in terms of writing style.

this approach may not work well for multiple scripts. Tan et al. (2016) proposed a method for selecting multi-features to define mutual information and used this for gender identification. Geometrical and transformed features are extracted from input handwritten document images. These features are effective for high quality images but less so for poor quality images impacted by degradations and distortion. Akbari et al. (2017) proposed a method based on a wavelet transform and probabilistic finite state automata for gender identification using handwritten text. The scope in this case is confined to a specific script and dataset. Sokic et al. (2012) proposed analysis of offline handwritten text samples using shape descriptors. The method extracts curvature, slope-based features, and Fourier descriptors for differentiating male and female handwritten text. However, the method requires individual character components for achieving its results. It is hard to segment the character components from Arabic and other Indian scripts without losing the shape of the characters. Liwicki et al. (2011) proposed automatic gender detection using online and offline information. The method explores a Gaussian Mixture Model, combining online and offline information for gender classification. Since the scope of our work here is gender identification from offline handwriting, their method may not apply to our data. Siddiqi et al. (2015) used automatic analysis of handwriting for gender identification. Their method extracts attributes including slant, curvature, texture, and legibility by using local and global features. Finally, an SVM classifier is used for classification. Since the method extracts language-specific features for gender classification, the performance will likely degrade for multi-script data.

Navya et al., (2018b) explored multi-directional Sobel kernels for feature extraction to identify gender using handwritten text lines. Their approach calculates correlation coefficients for successive text lines to identify consistent and inconsistent writing styles. If there is no change in the correlation coefficients for successive text lines, the process is said to be converging. Based on this criterion, their approach classifies a given handwritten input as produced by a male or female. The values of Sobel kernels are fixed in this work, which is a drawback of the method. To alleviate this problem, the same authors proposed a step for obtaining the values for the Sobel kernels (Navya et al., 2018b). This uses neighboring pixel information for calculating the values automatically. Then it uses the same steps stated in Navya et al., (2018a) for gender identification at the text line level. For applying their converging / diverging criteria, the approach requires at least three successive lines of text (Navya et al., 2018a; Navya et al., 2018b) and hence is not suitable for words. Ahmed et al. (2017) proposed handwriting-based gender classification using ensemble classifiers. Their method extracts texture features which include local binary pattern histograms of oriented features and statistical features computed from gray level co-occurrence matrices. In addition, the features are extracted through segmentation-based fractal texture analysis. For classification, the method uses SVM and ANN. Since the method extracts a large number of features, it is computationally expensive. In addition, it has not been tested on Indian scripts. Gattal et al. (2018) proposed gender classification from offline multi-script handwritten text using oriented basic image features, which consider the angle information of pixels. Their method uses a combination of different configurations of oriented basic image features. The histogram operation is performed for the extracted features to train an SVM classifier for gender classification. However, once again the above-mentioned approaches are applicable to text line images but not word images. Moetesum et al. (2018) proposed a method based on convolutional neural networks and linear discriminative analysis for gender identification at the text line and word levels. Their work addresses Arabic and Roman scripts, but the results are not consistent for text lines and words. Gattal et al. (2020) proposed a method for gender identification using COLD and hinge features for handwriting. The method uses cloud-of-line distribution for feature extraction from handwriting images. In addition, it captures curvature and traces the contour through hinge features. The feature matrix is supplied to a support vector machine for classification. The features are not robust to distortion, scaling,

and rotations, however. Kaljahi et al. (2019) used Gabor-based features for gender identification at the word level. For segmented words, weighted gradient features and clusters are used. Their approach performs a fusion operation for the histogram of Gabor responses for each word and then uses a CNN for gender identification. Since Gabor responses are sensitive to rotation and scaling, the approach is not invariant to these operations. In addition, the performance of the method depends on the fusion operation.

Overall, our review of methods for gender identification shows that most focus on foreground information, such as character shape, texture, and appearance. On the other hand, the methods ignore background information. This lack of information at the word level can lead to poor performance. It is also observed that none of these methods are independent of the script. In the case of multi-lingual countries, it is common to see single images with handwritten text indifferent languages. These observations motivate us to propose a method for addressing such challenges that arise in gender identification at the word level.

Our work aims at developing a new method based on Multi-Orientation-Scale Gabor Response Fusion (MOSGF). The style of writing includes particular usages of character shapes, contours, thickness of strokes, force, direction and speed of writing which are valuable distinguishing features (Navya et al., 2018a). To exploit this observation, the proposed method employs Gabor filters at different orientations and scales for enhancing fine details in handwriting (Ma et al., 2019). For each Gabor response, motivated by the method in (Ding and Wen, 2019) where correlation coefficients and Mahalanobis distance are used to capture global consistency and local similarity to understand radar images, our model adapts similar techniques in a novel way for gender identification. For the word segmentation problem, the approach presented in (Khare et al., 2018) for using weighted gradient features to segment touching handwritten text lines into words is adapted.

3. Proposed methodology

This section presents, at first, a method for segmenting words from handwritten text line images and then gender identification at the word level. Note that this work considers words that contain more than two characters for the gender identification application. As mentioned in the previous section, segmenting words from multi-script handwritten text lines is not easy. This is due to the cursive nature of characters in the case of non-Latin scripts, and the presence of calligraphic symbols and diacritics in the case of Arabic and Farsi scripts. As a result, conventional methods which use projection profiles may not work well (Moetesum et al., 2018). To address these challenges, begin by observing that the number of transitions from zero-to-one and one-to-zero (zero crossing points) at the top and bottom rows in text lines is less than that of zero crossing points at the middle rows for almost all of these scripts. This makes sense, because of the presence of ascenders, descenders, diacritics, and calligraphic symbols. Our idea, then, is to remove the pixels which represent the top and bottom portions of ascenders, descenders, diacritic and calligraphic symbols to facilitate identification of the spacing between words. The proposed method calculates the number of transitions (zero crossing points) for every row in the image and multiplies this value with the gradient of pixels in the respective rows. This results in a weighted gradient image, which contains high values for the middle rows and low values for the top and bottom rows. Then K-means clustering with $K = 2$ is performed on the weighed gradient image to separate the pixels that have high values from those that do not. This results in a Max cluster containing key components. The spaces between the components in this cluster are analyzed to determine the word segmentation. The flow of the logic for our word segmentation approach can be seen in Fig. 2.

Our expectation is that women usually write more neatly and maintain a consistency in their writing style and spacing in contrast to handwriting by method (Moetesum et al., 2018). It is also true that these characteristics are reflected in the orientation of contours and the

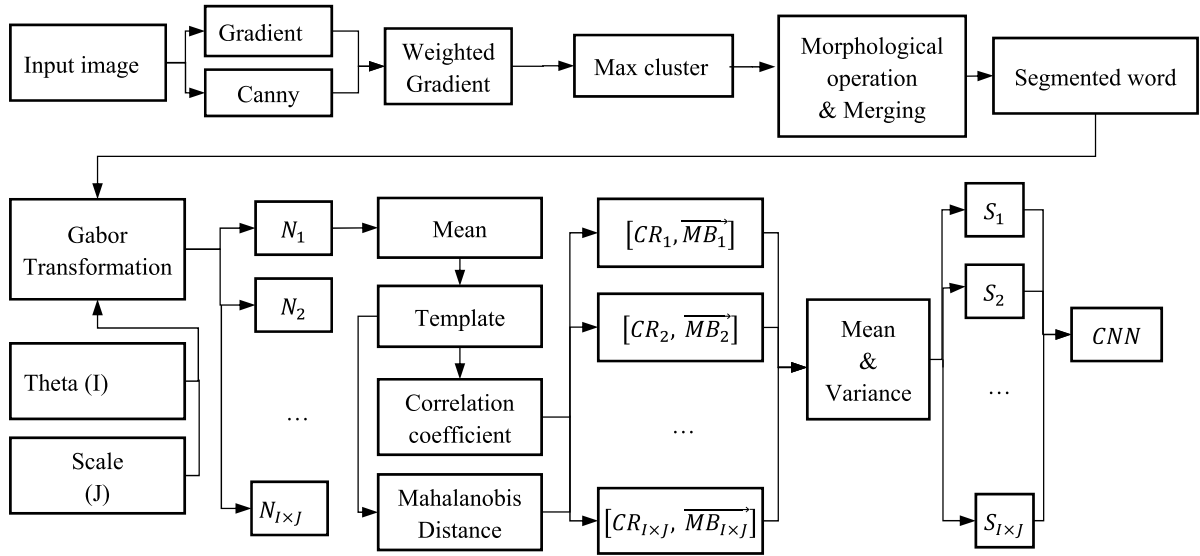


Fig.2. Framework of the proposed method for multi-script word segmentation.

overall structure of the writing. These observations motivate us to propose Multi-Orientation-Scale Gabor Responses (MOSGR) to extract differences in the orientations of the contours and the scale parameters (Ma et al., 2019). This process results in $6 \times 6 = 36$ Gabor responses for each word image, where 6 has been determined to be an optimal value for the theta and scale parameters of the Gabor function.

To study global consistency and local similarities in writing style, motivated by the method (Ding and Wen, 2019) where global and local information is used for recognizing radar images, the same general approach is adapted for gender identification at the word level. It can be expected that the spatial correlations between the Gabor responses and the average templates (the means of the 36 Gabor responses) will exhibit global consistency reflecting the structure of the handwriting. Moreover, the local similarity between the Gabor responses and the average template will highlight local changes in writing style. With this in mind, our proposed model estimates the spatial correlation coefficients and Mahalanobis distances between each Gabor response and the average template. In order to fuse the global and local features, the method calculate weights along with means and variances, which gives fused features for each input word image. For identification of female and male writing, the fused features are fed to a Neural Network (NN) (Arora and Suman, 2012) given their strong discriminating ability. The pipeline of the proposed work can be seen in Fig. 2, where S denotes the final fused features, N denotes the Gabor responses, and $[CR, \overrightarrow{MB}]$ denotes the feature vector given by the global and local methods.

3.1. Word segmentation

For handwritten text line images as illustrated in Fig. 3 (a), the proposed approach calculates the number of zero crossing points for each row using the Canny edge image of the input. The effect of the number of zero crossing points can be seen in Fig. 3 (b), where the number of zero crossing for the middle rows is higher than that for the top and bottom rows for both female and male writers. This is the advantage of using the number of zero crossing points, which is treated as the weight of the respective rows as defined in Equation (1) For edge pixels, gradients provide high values and low values of non-edge pixels (background). To take the advantage of this property to differentiate text and non-text pixels, the method multiplies the weight with the gradient values as defined in Equation (2), resulting in a weighted gradient image. The effect of the weighted gradient values is illustrated in Fig. 3(c), where one can see the values for middle rows are increased, while the values for top and bottom rows are decreased compared to the

values in Fig. 3 (b). This is also evident in Fig. 4 (a), where it is observed that the values for middle rows are enhanced compared to those for top and bottom rows.

Since the previous step widens the separation between middle row pixels and top-bottom row pixels, the proposed model performs K-means clustering with $K = 2$ to classify the larger values into a Max cluster and smaller values into a Min cluster, as defined in Equation (3) and Equation (4), respectively. Since the pixels classified into the Min cluster do not contribute to word segmentation they are discarded, and only the Max cluster results are considered in the next step. Fig. 4(c) shows that Max cluster contains almost all middle row pixels. As a result, the spacing between words is more prominent. At the same time, the removal of Min cluster pixels may lead to disconnections. Therefore, the method performs a morphological operation to close the gap between pixels as shown in Fig. 4(d), which outputs word components by merging sub-components. The proposed system fixes bounding box for each component as shown in Fig. 4(e). In the event the bounding box of a component misses any sub-components due to a large space, the method checks the area of each bounding boxes. If overlapping regions are found between two bounding boxes, the method merges the two bounding boxes as one, as defined in Equation (5). The final word segmentation can be seen in Fig. 4(f). The word segmentation results illustrated in Fig. 5 show that the proposed approach works well irrespective of script and orientation.

$$Z_r = \sum_{l=1}^U C_{x,y} \quad (1)$$

where Z_r is the weight of the r^{th} row, U is the number of columns in the canny image, C and $C_{x,y} = \{0,1\}$.

$$ZG_{x,y} = Z_x \times G_{x,y} \quad (2)$$

where $ZG_{x,y}$ is the weighted gradient of each pixel (x,y) in the weighted gradient image, and $G_{x,y}$ is gradient-magnitude of each pixel.

Then, our method clusters the ZG images into two classes:

$$Max_{x,y} = \begin{cases} ZG_{x,y} & \text{if } \text{means}(ZG_{x,y}) \in cluster_{max} \\ 0 & \text{else} \end{cases} \quad (3)$$

$$Min_{x,y} = \begin{cases} ZG_{x,y} & \text{if } \text{means}(ZG_{x,y}) \in cluster_{min} \\ 0 & \text{else} \end{cases} \quad (4)$$

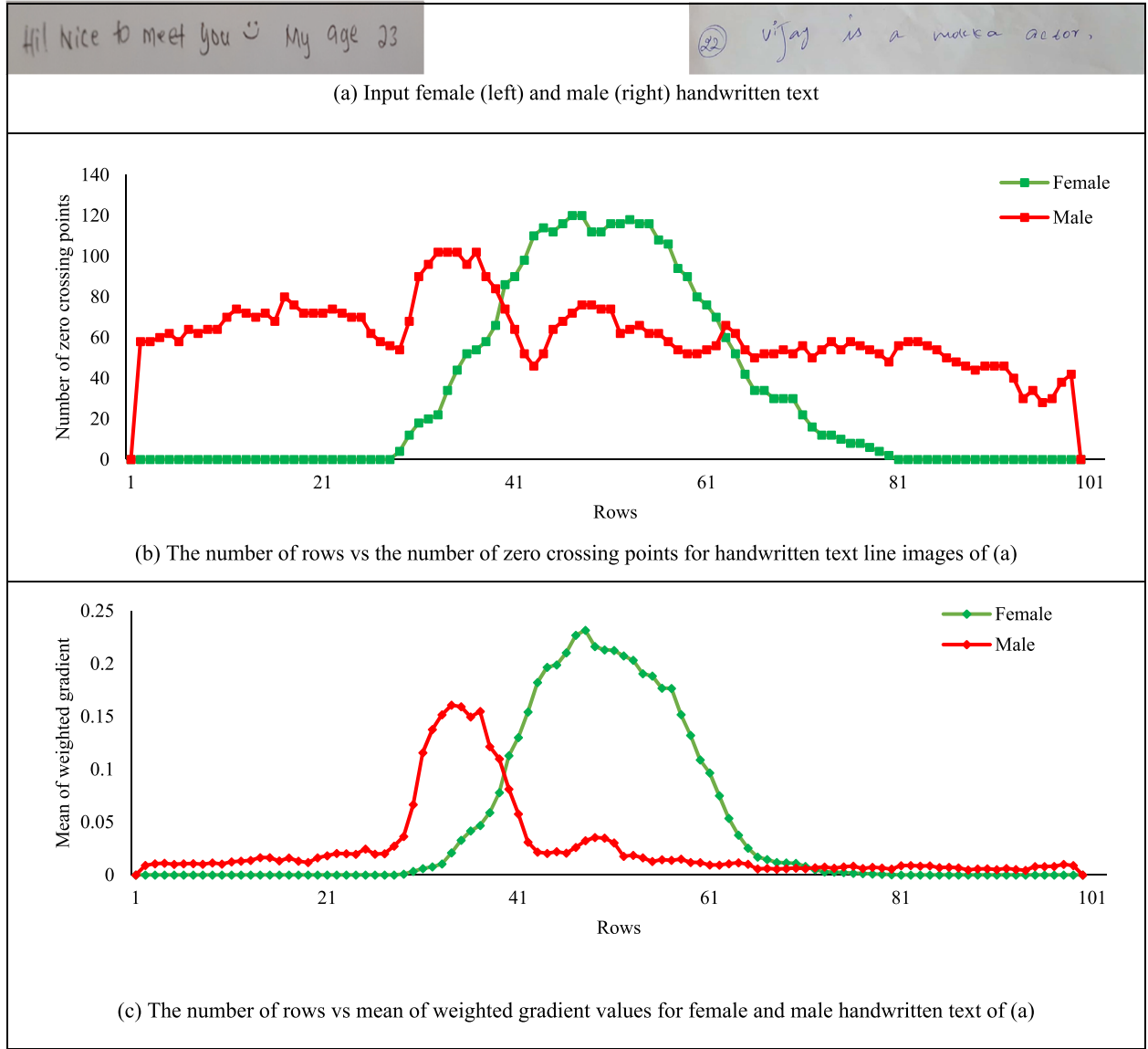


Fig.3. The combination of zero crossing points and gradient values for classifying middle and bottom-top pixels of text lines.

$$B_i = \begin{cases} 1 & \text{if } \text{intersection}(\text{Box}_{Max}^k \text{ and } \text{Box}_{Max}^g) = \text{True} \\ 0 & \text{else} \end{cases} \quad (5)$$

where, in Equation (5), t indicates the number of bounding boxes in the Max cluster, (k, g) denotes peer bounding boxes that overlap the bounding box being checked. $B_i = 1$ means the proposed method merges the two bounding boxes, otherwise it them separate in the final word segmentation result.

3.2. Multi-Orientation-Scale Gabor responses (MOSGR) for enhancing fine details

For each word segmented in the previous step, as mentioned earlier, the proposed system applies Gabor filters, as defined in Equation (6), to obtain Gabor responses for different orientations and scales, as shown in Fig. 6. In the spatial domain, a two-dimensional Gabor filter is a Gaussian kernel function modulated by a complex sinusoidal plane wave as defined in Equation (6) (Ma et al., 2019). As defined in Equation (7), the proposed method generates 36 Gabor responses by considering orientations and scales. The outputs of the Gabor filters enhance fine details in handwritten words as shown in Fig. 6. In order to extract the

global consistency and local similarity at the word level, as discussed earlier, our method generates template by taking the averages of all the Gabor responses as defined in Equation (8). The 36 Gabor responses, which include 6 for theta and 6 for scale, are shown in Fig. 6. The values for the theta and scale parameters are determined empirically, as described later in the Experimental Section.

$$D(x, y, \delta, \gamma) = \frac{d^2}{\pi\alpha\beta} \exp\left(-\frac{x'^2 + \alpha^2 y'^2}{2\gamma^2}\right) \exp(i2\pi r x' + \mu) \quad (6)$$

$$x' = x \cos \delta + y \sin \delta$$

$$y' = x \sin \delta + y \cos \delta$$

where d represents the frequency of the sinusoid, δ is the orientation of the normal to parallel stripes of a Gabor function, μ is the phase offset, γ is the standard deviation of the Gaussian envelope, and α is the spatial aspect ratio which specifies the ellipticity of the support of Gabor function. More details about Gabor filters can be found in (Ma et al., 2019).

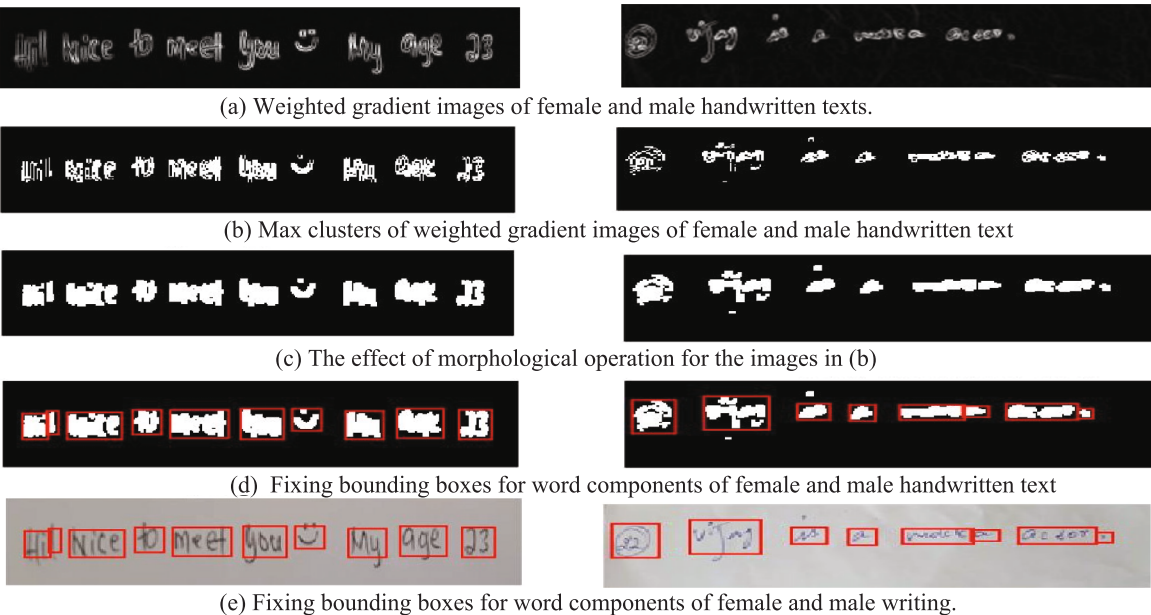


Fig.4. The proposed word segmentation for the female and male handwritten text lines.

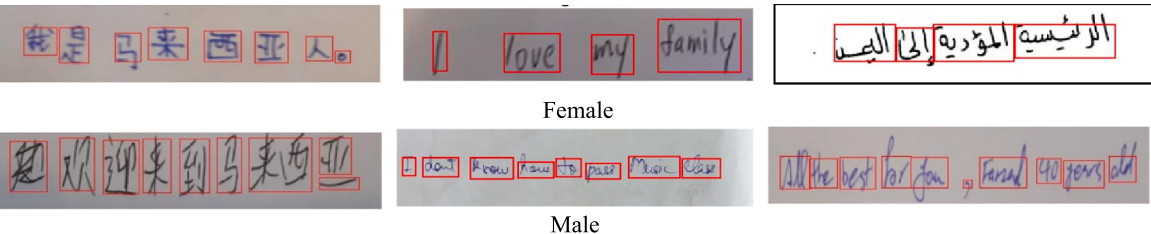


Fig.5. Sample of proposed word segmentation results for the different scripts with alignments.

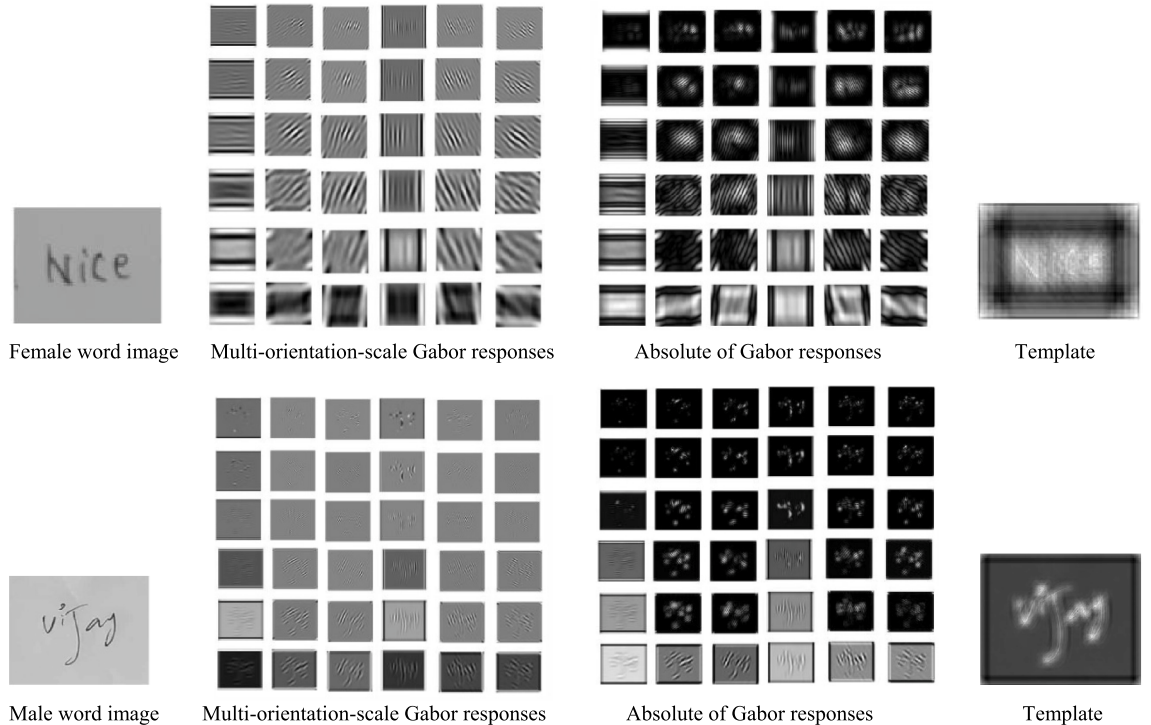


Fig.6. Multi-orientation and scale of Gabor responses for enhancing the fine details in female and male word images.

$$N(\delta_i, \gamma_i) = G(x, y) \times D(\delta_i, \gamma_i) \quad (7)$$

where $\delta_i, i = \{1, 2, \dots, I\}$ and $\gamma_j, j = \{1, 2, \dots, J\}$, N denotes the set of Gabor response, G is a grayscale image and D is the Gabor filter bank.

$$T = \frac{\sum_{i=1}^I \sum_{j=1}^J N(x, y, \delta_i, \gamma_j)}{I \times J} \quad (8)$$

3.3. Fusing correlation and distances based features for gender identification

Since the target of our work is to study the global consistency and local similarity between individual Gabor responses and templates for extracting distinctive features, spatial correlation and local similarities are estimated using correlation coefficients and Mahalanobis distance as defined in Equation (9) and Equation (10), respectively (Ding and Wen, 2019). Correlation Coefficient (CR) is computed for each Gabor response and template, which outputs a feature vector containing 36 values. Since CR is obtained by considering the whole image response and template, it provides global information. The effect of CR is illustrated in Fig. 7(a), where it is noted that for the 36 Gabor responses, CR gives visually different values for female and male handwriting. Similarly, for a row of a Gabor response and the corresponding row of the template, the proposed method computes similarity using Mahalanobis Distance (MB), which results in a feature matrix with the same number of rows as the input image. The effect of MB can be seen in Fig. 7(b), where it is observed that the graphs of female and male writers are different. Note: for the Y axis in the graph, the method considers the mean of MB of all the rows. From Fig. 7(a) and Fig. 7(b), it can be concluded that CR and MB are useful in extracting distinctive aspects of female and male handwritten text. The feature vectors CR and MB can be combined as a vector $\vec{V} = [CRMB_1 MB_2 \dots MB_R]$, where R denotes the total number of rows in the image. To fuse these two feature vectors, the proposed method normalizes the values by assigning the weights as defined in Equation (11). Equation (11) ensures the sum of the weights should be equal to 1. Therefore, the method assigns $w_0 = 0.5$, which is half of the total weight to the CR feature vector and the remaining half to the MB feature vector. This results in the feature vector as defined in Equation (12). The final fusion score is extracted by computing the mean and variance as defined in Equation (13), which gives a vector of 36 values. The distribution of fusion scores is plotted in Fig. 7(c), where it can be seen that the graphs are smoother than those in Fig. 7 (a) and Fig. 7 (b). In this way, the proposed method integrates global consistency and local similarity to create distinctive features for gender identification.

$$CR(\delta_i, \gamma_i) = \frac{\sum \sum (T - \bar{T}) \times (N(\delta_i, \gamma_i) - \bar{N}(\delta_i, \gamma_i))}{\sqrt{(\sum \sum (T - \bar{T})^2) \times (\sum \sum (N(\delta_i, \gamma_i) - \bar{N}(\delta_i, \gamma_i))^2)}} \quad (9)$$

Where T denotes template, N denotes the Gabor responses, $\bar{T}$ and $\bar{N}(\delta_i, \gamma_i)$ are the mean values of T and N , respectively. $CR(\delta_i, \gamma_i)$ is the correlation coefficient between template and response with δ_i, γ_i scale and orientation (theta), respectively.

$$MB(\delta_i, \gamma_i)^r = \sqrt{\left(\frac{N(\delta_i, \gamma_i)^r - \bar{T}^r}{\sum \sum (N(\delta_i, \gamma_i)^r - \bar{T}^r)^2} \right)} \quad (10)$$

where $\overline{N(\delta_i, \gamma_i)^r}$ represents the vector of the r^{th} row of response image of $N(\delta_i, \gamma_i)$ if $r = \{1, 2, \dots, R\}$. R is the total number of the rows in the image. T^r denotes corresponding rows in the template image, and $MB(\delta_i, \gamma_i)^r$ is the distance for the r^{th} row of $N(\delta_i, \gamma_i)$ to the template.

$$\begin{cases} w_0 + w_1 + \dots + w_R = 1 \\ w_0 = 0.5 \end{cases} \quad (11)$$

$$\vec{WV} = \vec{V} \times \vec{W} \quad (12)$$

where $\vec{V}$ is multiplied by weight vector.

$$S(\delta_i, \gamma_i) = \omega \times \exp\left(\frac{-\theta^2}{2}\right) \quad (13)$$

where $S(\delta_i, \gamma_i)$ denotes Score of Gabor responses, which is considered as features, ω and θ^2 denote mean and variance of $\vec{WV}(\delta_i, \gamma_i)$, respectively

The extracted features are supplied to a fully connected NN for classifying female and male handwritten text at the word level (McAllister et al., 2016). Here the paradigm described in (Nanni et al., 2018) is that as the number of training samples increases, the performance of the NN also increases. While bigger training sets produce better results for NN classifiers, creating and labeling a large database for gender identification is challenging. Hence, our approach is to use a combination of the features just described rather than raw pixels as with recent deep learning models. Since our main focus is developing powerful new features for gender classification from handwriting, for classification the popular WEKA platform (Arora and Suman, 2012) is used, employing its built-in NN and using the default parameter settings and architecture. The NN has 5 fully connected layers (FC layers) with 36 neurons in the input layer and 1 neuron in the final output layer. The Binary Cross Entropy loss function (Nasr et al., 2002) is used for training, the “ReLU” activation function in the hidden layers, and “Sigmoid” (Narayan, 1997) as the activation function for the final layer. The Adam optimizer (Kingma and Bai, 2015) with a learning rate of 0.005 is used during the training process, with the batch size of 8 for 50 epochs. The model architecture is shown in Fig. 8. This is a binary classifier which outputs ‘0’ for the female class and ‘1’ for the male class.

4. Experimental results

To the best of our knowledge, this is the first work on gender identification at the word level using as input handwritten text line images in a multi-script environment, including Arabic, Farsi, Chinese, and several Indian scripts. For this reason, a new dataset has been created for experimental purposes. This dataset consists of handwritten words images reflecting different styles, orientations, backgrounds, papers, inks, and writer ages ranging from 10 to 70 years old. In total, our dataset contains 2,278 text lines from female writers, and 1,026 text lines from male writers. For comparison purposes, the proposed method is also tested on standard benchmark datasets, namely, IAM which contains English text lines, KHATT which contains Arabic text lines, and QUWI which contains both English and Arabic text lines. The aforementioned datasets comprise 100 text lines per class, 45 text line per class, and 250 text lines per class, respectively. In summary, a total 4094 text line images are used for the experiments (3,304 + 200 + 90 + 500). Note that this new dataset is larger and more varied compared to the standard datasets.

Sample images for female and male writers for the various datasets are shown in Fig. 9.

In this work, word segmentation from handwritten text line images and gender identification at the word level are the key steps. To evaluate the proposed approach to word segmentation, this work employs the standard measures mentioned in (Khare et al., 2018), namely, Recall (R), Precision (P) and F-Measure (F). Since the ground truth for our dataset is not available, the measures are counted manually. For gender identification, the approach of (Moetesum et al., 2018) is followed, where classification rate is used for evaluation. As noted earlier, in this work, words are assumed to contain more than two characters. The total numbers of words in the datasets used for the evaluation are reported in Table 1.

To show effectiveness of the proposed approach for word segmentation, comparisons are made to existing methods (Louloudis et al. 2016; Banumathi et al., 2016; Lamsaf et al., 2019) which are based on a connected component analysis. The method in (Banumathi et al., 2016)

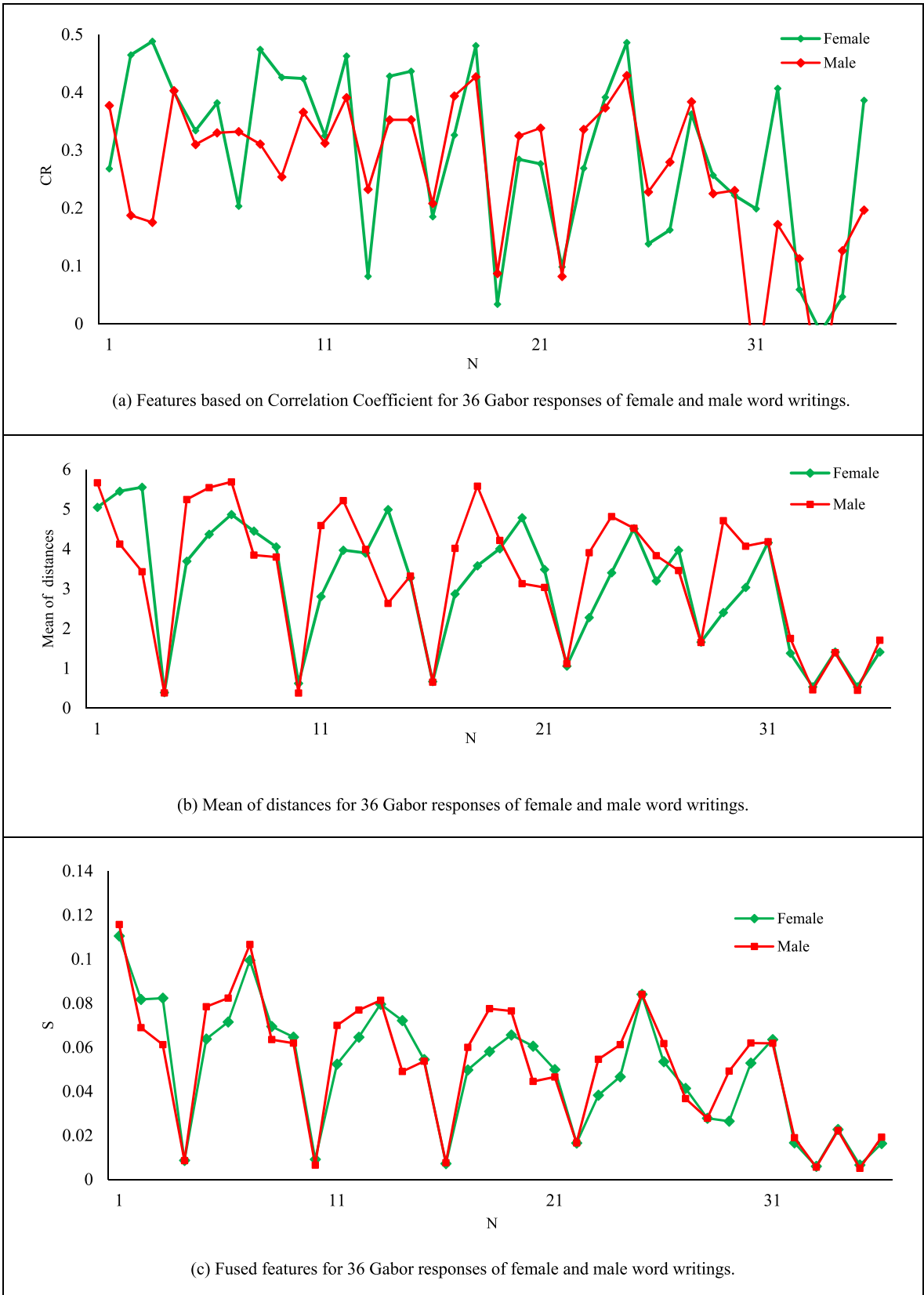
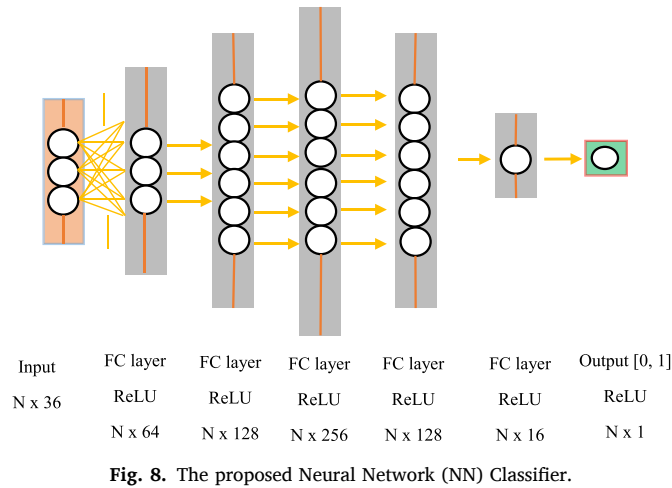


Fig.7. The combined features for 36 Gabor responses of female and male word writings. N denotes the number of Gabor responses and S denotes final fused features.



aims to segment words in Indian script, while the method (Louloudis et al. 2016) aims to segment words in historical handwritten document images. The method (Lamsaf et al., 2019) aims to segment words in Arabic handwritten document images.

For gender identification, the following state-of-the-art methods are considered for a comparative study. Bouadjene et al. (2016) describes an approach that uses a fuzzy concept for gender identification, and Bouadjene et al. (2015) presents a method that uses descriptors and an SVM classifier for gender identification. Navya et al., (2018a) and Navya et al., (2018b) uses multi-gradient Sobel kernels for gender identification. The motivation to choose these methods for comparative study is that they use concepts similar to our method. However, the methods (Navya et al., 2018a; Navya et al., 2018b) require at least three handwritten text lines for successful gender classification. For a fair comparative study, these gender classification experiments are done at the text line level. However, the other two methods, (Bouadjene et al., 2016; Bouadjene et al., 2015), work at both the line and word levels. Gattal et al. (2020) presents a method using a cloud of line distributions and hinge features for gender identification. Kaljahi et al. (2019) describe an approach which extracts the Gabor responses using a Gabor transform for gender identification at the word level. This comparative study examines the robustness of these methods relative to the approach being proposed.

In the method, the values for theta and the scale parameters in the

Gabor function have a significant impact on performance. Therefore, the optimal values for those two parameters were determined using our dataset as shown in Fig. 10, where accuracy is calculated for varying theta (I) and fixing the scale parameter (J). Similarly, in the second graph in Fig. 10, accuracy is calculated for varying the scale (J) and fixing the theta parameter. It is observed from the first and the second graphs in Fig. 10 that the accuracy increases gradually as the value of the parameter increases. The accuracy at the value 6 provides the maximum in both cases. Therefore, the value of both theta and the scale parameter is set to 6 in these experiments.

4.1. Evaluation of the word segmentation method

Sample image results for word segmentation are shown in Fig. 11, where it can be seen that the proposed approach segments words successfully for different scripts and different alignments. It is also observed from Fig. 11 that the proposed approach works well for images affected by low contrast, low resolution, and to some extent to blur and distortion. This also demonstrates that the method is effective for word segmentation in multi-script and multi-orientations scenarios. A quantitative comparison of the method and the existing approach on the different datasets is reported in Table 2, where the proposed approach has the best F-measure compared to the existing approaches for all of the datasets. As noted previously, the existing methods were developed with a focus on particular scripts and the extracted features are not adequate to handle the challenges introduced by calligraphic symbols in the case of Arabic and Farsi. However, the Lamsaf et al. (2019) method yields a better F-measure compared to the other two existing methods. This is because it was developed to handle Arabic word segmentation. Our approach is robust to the challenges introduced by calligraphic symbols, ascenders / descenders due to the features used. Moreover, the method produces nearly consistent results across three of the four datasets.

Since the proposed approach should be robust to different rotations, scaling, and to some extent noise, images from the new dataset were

Table 1

The number of handwritten text line images of different datasets considered in this work for gender classification at word level.

Datasets	Proposed	IAM	KHATT	QUWI
Female	10,379	1036	714	1761
Male	8745	991	848	1595
Total	19,124	2027	1562	3356

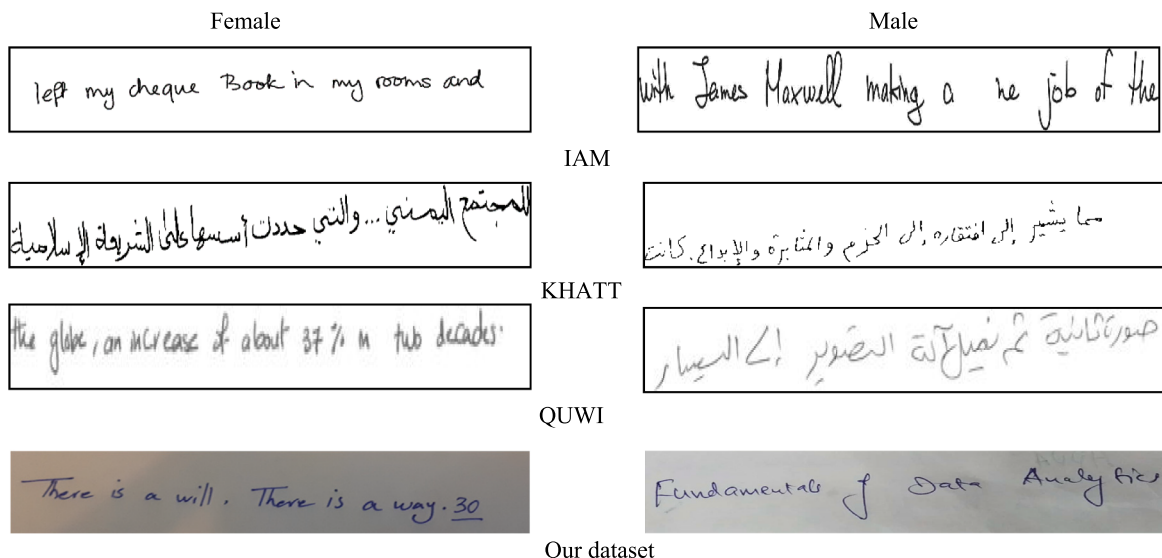


Fig.9. Some sample handwriting images of different databases.

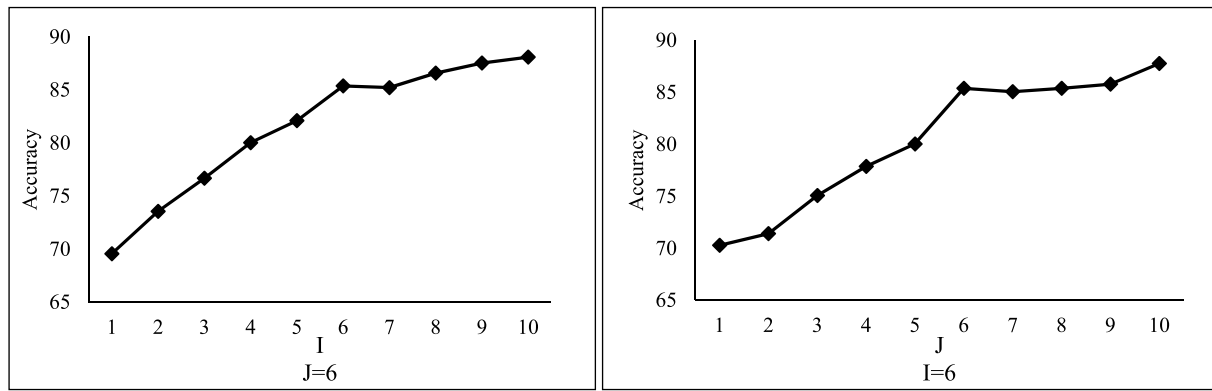


Fig.10. Determining the best number of orientations and scales of Gabor responses for handwriting word-based gender identification.

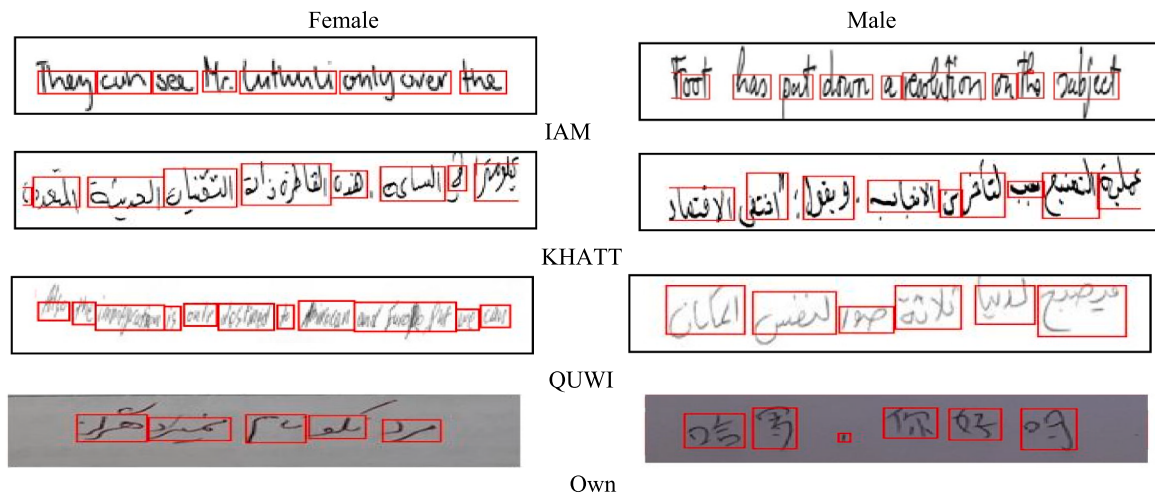


Fig.11. The results of word segmentation of our method for multiple scripts and orientated text lines.

Table 2

Measures of the proposed and existing approaches on our and standard datasets for word segmentation.

Methods and Dataset	Proposed Method			Banumathi et al. (2016)			Louloudis et al. (2009)			Lamsaf et al. (2019)		
	R	P	F-Measure	R	P	F-Measure	R	P	F-Measure	R	P	F-Measure
Our dataset	82.14	76.25	79.23	14.55	40.0	21.33	28.18	62.0	38.75	56.0	93.3	69.9
IAM	81.40	58.88	78.06	55.99	77.65	65.06	70.36	46.19	55.77	70.9	85.6	77.5
KHATT	55.55	69.19	61.63	10.37	56.00	17.50	30.36	47.71	37.11	43.0	77.8	51.5
QUWI	78.60	79.45	79.03	14.66	47.16	22.35	44.58	37.65	40.82	73.0	76.7	73.5

rotated randomly to further explore the robustness. Similarly, images were scaled up and down randomly by 0.2–3 times their original size. Random Gaussian noise was also added, ranging from a level of 0.001–0.1. Sample results of word segmentations under these conditions are shown in Fig. 12 where it can again be seen that the proposed approach segments the words successfully for the different rotations, various scaling, and noise level. Comparative results between this new approach and the existing methods are reported in Table 3, where it is noted that our approach yields the best results by a significant margin. This makes sense because the steps used, namely, candidate pixels weighted by gradient values and clustered using k-means, reduces distortions that result from rotations, scaling, and added noise. At the same time, the nearest neighbor criterion used for merging character into words does not depend on rotation and scale. Although the proposed approach counts the number of zero-crossing points horizontally, this has no effect for different rotations when considering the overall performance due to the use of gradient values and k-means clustering. However, since existing methods are not invariant to rotation, scaling,

and noise, they suffer greater impacts from these common distortions. When the results of the proposed and existing methods are compared in Table 2 and Table 3, all methods score lower for the distorted dataset. It can be noted that in the case of the proposed approach, the difference is less than that of the other methods. Note that since the (Lamsaf et al., 2019) method is sensitive to distortion, noise, and blur, the method does not work in such cases and hence its results are not reported in Table 3.

4.2. Ablation study

The proposed method for gender identification involves parameter tuning to achieve the best possible results, namely, the scale and theta (orientation) parameters in the Gabor function, the Correlation Coefficients (CR) based features for global consistency extraction, and the Mahalanobis (MB) Distance measure for local similarity estimation. To assess settings for these parameters, a series of experiments were conducted on our dataset with the results reported in Table 4. When the results are compared with only scale without varying orientation, and

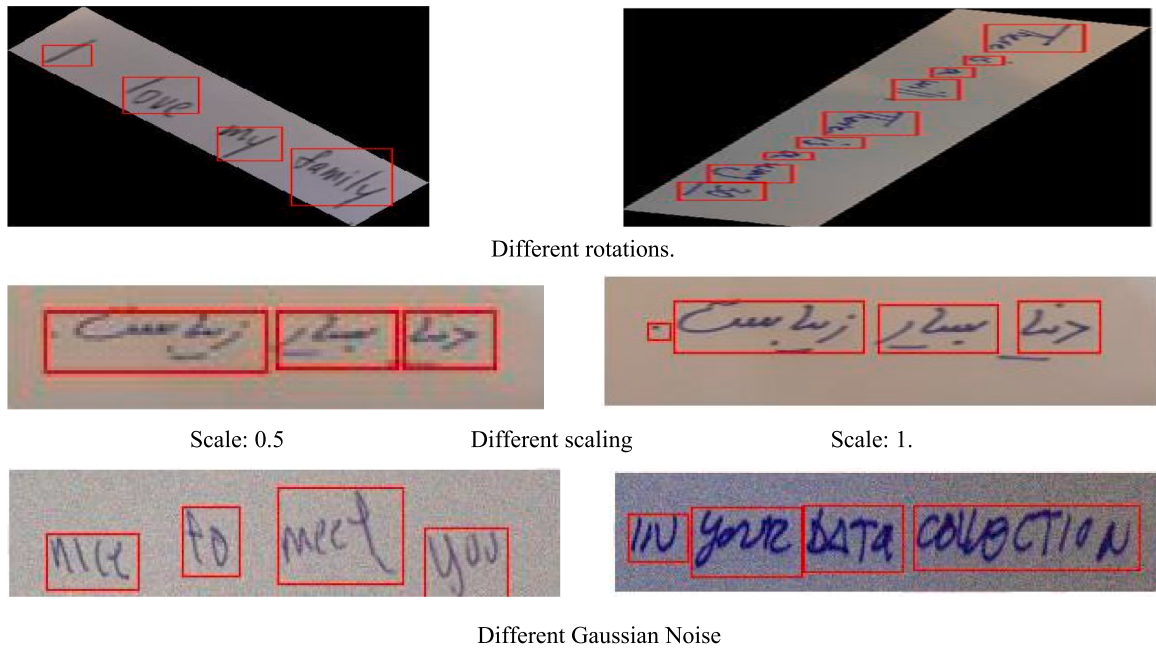


Fig.12. The proposed word segmentation for the image of different rotations, scaling and noise.

Table 3

Testing the proposed and existing methods robustness on different rotations, scaling and injected noise images from our dataset.

Distortion on own dataset	Proposed Method			Banumathi et al. (2016)			Louloudis et al. (2009)		
	R	P	F-Measure	R	P	F-Measure	R	P	F-Measure
Different Rotations	74.55	63.01	68.30	0.08	0.3	12.85	23.83	41.80	30.35
Different Scales	62.71	74.49	68.09	21.76	50.68	30.45	37.64	43.83	40.50
Different Levels of Gaussian Noise	67.2	58.33	62.45	0	0	0	0	0	0

the results with only orientation without varying scale, it can be noted that both yield almost the same results. Hence both contribute equally when it comes to achieving gender identification results. Similarly, when results are compared using only CR without MB, and vice versa, the results are also fairly close. This allows us to conclude that both CR and MB are useful in achieving the best results in the proposed approach.

4.3. Evaluation of gender identification at the word level

Sample image results for gender identification on the different datasets are shown in Fig. 13, where one can see the proposed method successfully classifies the words of different scripts, alignments, orientations, and to some extent to blur and distortions. Quantitative results comparing the proposed and existing methods on the different datasets are reported in Table 5, where both confusion matrices and classification rates can be seen. The classification rate of the proposed approach is better than those of the existing methods for all the datasets. This is due to the attention paid to identifying gender at the word level as opposed to the text line level, as well as to the robustness of the approach to multi-script inputs. The method described in (Gattal et al., 2020) yields the highest classification rate among the other methods for most of the datasets. However, it is still dominated by our approach for all of the

datasets. From Table 5, it can be seen that the existing methods score the highest for the IAM dataset. This is undoubtedly because the IAM dataset is limited to English text lines. When the results of all the datasets including our dataset are compared, the proposed approach scores consistent results while the existing methods exhibit more variability. To visualize the performance of the proposed method with respect to the existing approaches, line graphs for the different datasets are plotted in Fig. 14. The similar conclusions observed from the Table 5 can be drawn from the graphs shown in Fig. 14.

As described in the previous section, all the methods can also be tested for different rotations, scaling, and noise levels by injecting these distortions into our dataset. Qualitative results for our approach are shown in Fig. 15, where the images are classified successfully in spite of the various distortions. This illustrates that the fusion step used for combining global and local features is robust. Quantitative results comparing the new approach to the previous methods are reported in Table 6, again showing that our approach achieves the best classification rates. Comparing the results of Table 5 and Table 6, it can be seen that the proposed approach is largely robust to geometrical transformations for multiple script inputs.

To visualize the performance of various methods across the different distortion types, line graphs are plotted in Fig. 16. The same conclusions

Table 4

Analyzing the contribution of the key steps of the proposed gender identification method (F: Female, M: Male).

Measures	(CR + MB) Scale		(CR + MB) Orientation		CR		MB		Proposed method	
	F	M	F	M	F	M	F	M	F	M
F	72.99	27.0	73.07	26.92	77.81	22.18	74.75	25.24	78.52	21.47
M	40.28	59.71	39.15	60.84	32.94	67.05	24.5	75.49	23.05	76.94
Classification Rate	66.35		66.95		72.43		75.12		77.73	

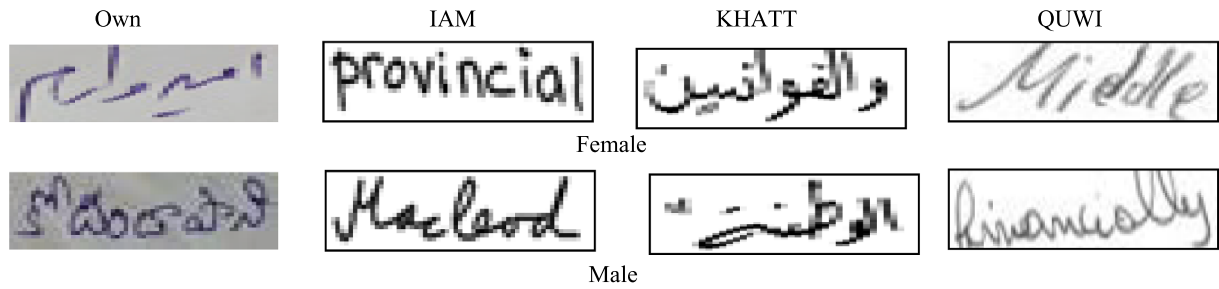


Fig.13. Sample of the successful multi-script and multi-orientations gender identification results at word level of the proposed method.

Table 5

Confusion matrix of the proposed and existing methods for gender identification at word level.

Methods	Measures	Datasets		IAM	M	KHATT	M	QUWI	M
		Our	F	F		F		F	
Proposed method	Female (F)	78.5	21.4	82.7	17.2	85.9	14.0	76.5	23.4
	Male (M)	23.0	76.9	19.6	80.3	18.6	81.3	20.3	79.6
	Classification Rate	77.7		81.5		83.6		73.6	
Bouadjenek et al. (2016)	Female (F)	58.9	41.0	80.4	19.5	48.6	51.3	47.6	52.3
	Male (M)	27.4	72.5	25.4	74.5	49.5	50.4	50.6	49.4
	Classification Rate	65.7		77.4		49.5		48.5	
Bouadjenek et al. (2015)	Female (F)	50.0	49.9	57.3	42.6	49.0	50.9	56.5	43.4
	Male (M)	47.4	52.5	45.0	54.9	42.4	57.5	54.5	45.4
	Classification Rate	51.3		56.1		53.3		51.0	
Gattal et al. (2020)	Female (F)	59.5	40.4	36.8	63.1	39.6	60.3	40.4	59.5
	Male (M)	9.2	90.7	17.5	82.4	25.6	74.4	23.8	76.1
	Classification Rate	75.1		59.6		57.0		58.3	
Kaljahi et al. (2019)	Female (F)	99.8	0.1	51.1	48.8	17.7	82.2	62.1	37.8
	Male (M)	79.4	20.5	39.1	60.8	13.3	86.6	46.8	53.1
	Classification Rate	60.1		56.0		52.1		57.6	

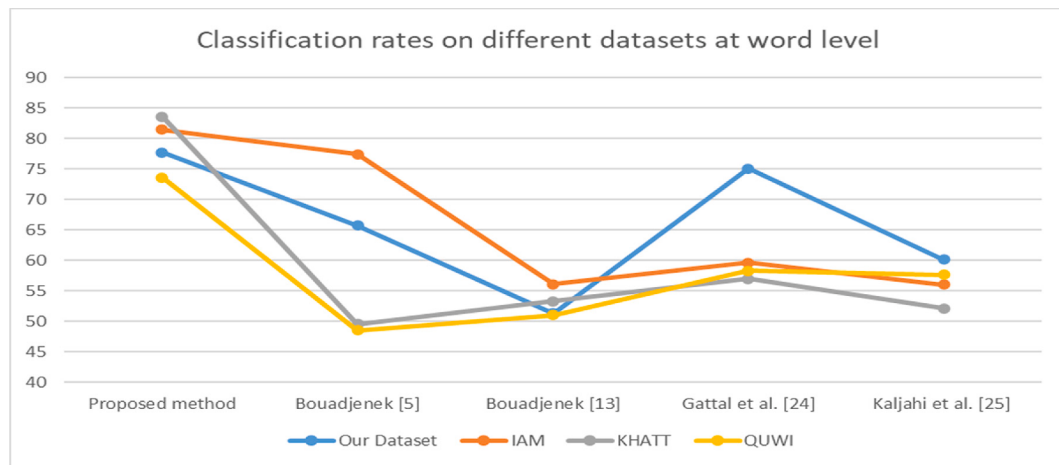


Fig. 14. Analyzing the results of the proposed and existing methods on different datasets for classification at word level.

that observed from Table 6 can be drawn from the graphs shown in Fig. 16.

4.4. Evaluating gender identification at the text line level

It can be expected that when a method works well at the word level, it should work at least as well at the text line level since more information will be present. To test this for the proposed approach and the existing methods, experiments were conducted at the text line level and report those results in Table 7. It can be noted by comparing Table 5 and Table 7 that there is a significant improvement when processing lines as opposed to words for all of the methods. And, once again, the proposed method dominates the other approaches across all of the datasets. For

this experiment, (Navya et al., 2018a; Navya et al. 1018b) method was included in the comparison because it requires a minimum of three lines for gender identification. To visualize the performance of the various methods on gender classification at the text line level across the different datasets, line graphs are presented in Fig. 17. It can be noted from the figure that the proposed approach outperforms the existing methods in all cases.

Although the proposed method is relatively robust for word segmentation, it sometimes fails for Arabic, Farsi, and other scripts as shown in Fig. 18. This is due to large variations in writing calligraphic symbols and diacritics in the case of Arabic and Farsi, and the use of ascenders and descenders in different styles for Roman (English) script, and compound characters in case of Indian scripts. Addressing such

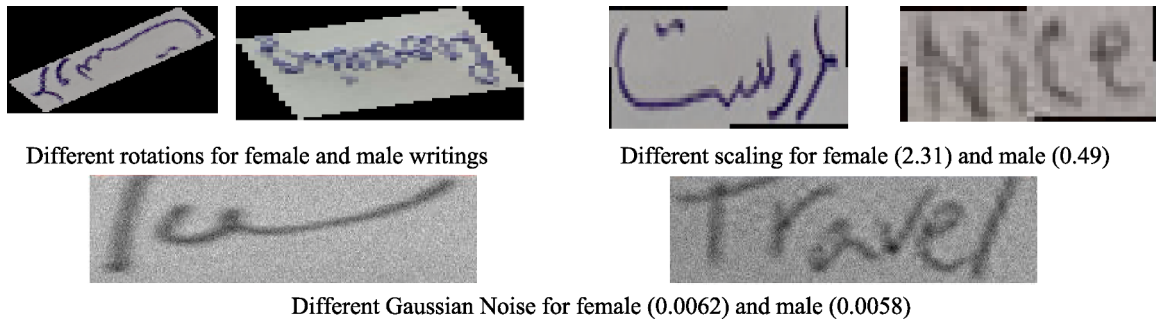


Fig.15. The proposed gender identification for the different rotations, scaling, and noisy images.

Table 6

Evaluating the proposed and the existing approaches on the images of different rotations, scaling, and noise for gender identification at word level.

Methods	Measures	Rotation		Scale		Noise	
		F	M	F	M	F	M
Proposed method	F	74.09	25.90	77.19	22.80	75.23	24.76
	M	39.47	60.52	31.15	68.84	38.28	61.71
	Classification Rate	67.88		73.37		69.05	
Bouadjene et al. (2016)	F	51.8	48.2	57.6	42.4	47.6	52.4
	M	48.8	51.2	39.4	60.6	47.2	52.8
	Classification Rate	51.5		59.1		50.2	
Bouadjene et al. (2015)	F	40.2	59.80	53.4	46.6	29.2	70.8
	M	39.2	60.8	40.2	58.81	25.4	76.5
	Classification Rate	50.5		50.6		51.9	
Gattal et al. (2020)	F	38.46	61.54	45.94	54.06	44.51	55.49
	M	29.07	70.93	24.46	75.54	24.22	75.78
	Classification Rate	54.7		60.74		60.14	
Kaljahi et al. (2019)	F	99.82	0.18	67.78	32.22	75.45	24.55
	M	81.01	18.99	53.73	46.27	66.42	33.58
	Classification Rate	59.4		57.02		54.52	

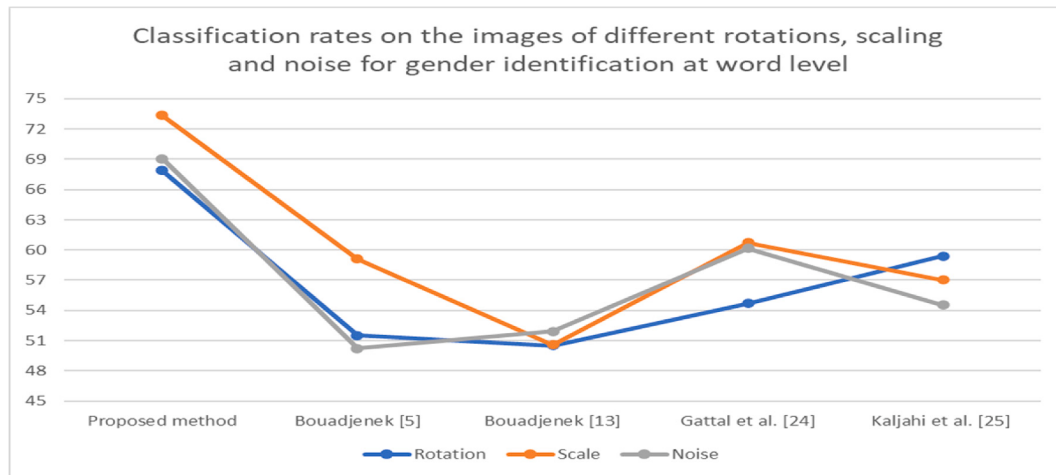


Fig. 16. Analyzing the results of the proposed and existing methods on different rotations, scaling, and noise of different datasets for classification at word level.

situations may require full-scale word recognition, which is a topic beyond the current paper. Of course, the approach will also fail for individuals who write in a way that is more similar in style to those of the other gender, as shown in Fig. 19. Therefore, there is still room for future improvement.

5. Conclusion and future work

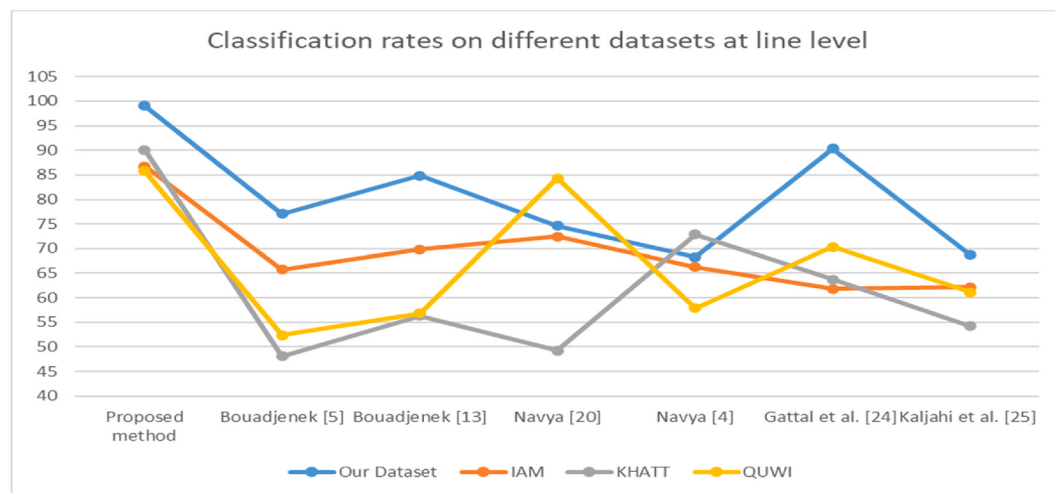
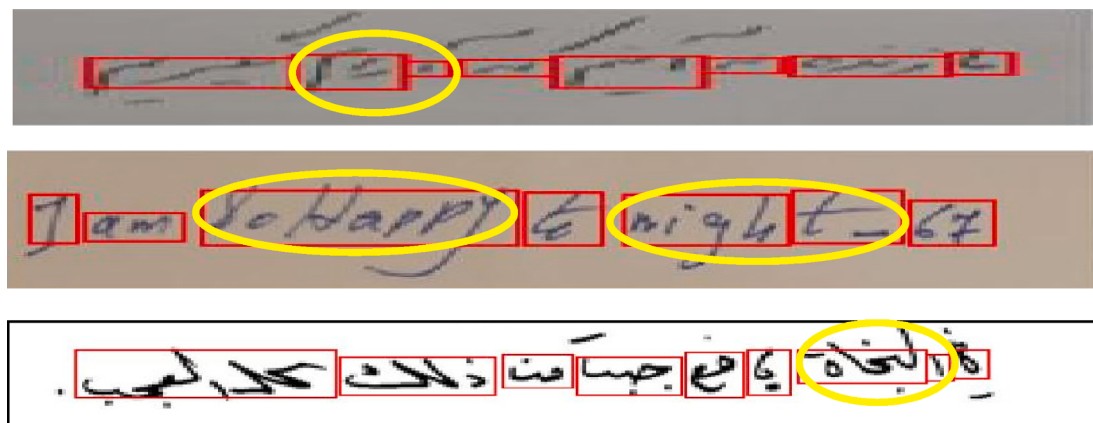
This paper has introduced a new system for gender identification at the word level based on handwritten text analysis. The approach uses weighted gradients for segmenting words in handwritten text line images, based on the number of zero crossing points. To extract distinctive

features at the word level, Gabor responses for different orientations and scales are employed. For each Gabor response for an input word image, correlation coefficients for global consistency and Mahalanobis distance for local similarity are extracted. The proposed method uses a new way to fuse the features based on weight determination. This results in a feature vector for the input image which is then fed to neural network for the classification of female and male writing. Experimental results using our own dataset and three standard datasets from the field demonstrate that our approach outperforms existing methods and exhibits substantial robustness. However, as noted in our experimental evaluation, word segmentation can still be challenging in certain cases, and various forms of variability and degradation can negatively impact

Table 7

Evaluating the proposed and existing approaches on different datasets for gender identification at line level.

Methods	Measures	Our F	M	IAM F	M	KHATT F	M	QUWI F	M
Proposed method	F	99.53	4.62	87.18	12.81	96.49	3.50	85.6	14.4
	M	1.92	98.07	13.57	86.42	12.59	87.40	14	86
	Classification Rate	99.04		86.80		90.12		85.8	
Bouadjenek et al. (2016)	F	64.23	35.76	66.07	33.92	23.56	76.44	55.2	44.8
	M	10.07	89.92	34.64	65.35	6.67	93.32	50.4	49.6
	Classification Rate	77.07		65.71		48.14		52.4	
Bouadjenek et al. (2015)	F	67.44	23.55	79.2	20.8	54.07	45.92	58	42
	M	7.07	92.92	34.8	65.2	41.48	58.51	44.4	55.6
	Classification Rate	84.88		69.82		56.29		56.8	
Navya et al., (2018b)	F	74.27	25.73	72.41	72.41	57.14	57.14	85.71	14.29
	M	25	75	27.50	72.50	58.62	41.38	17.03	82.97
	Classification Rate	74.63		72.45		49.26		84.34	
Navya et al., (2018a)	F	68.30	31.70	59.52	40.48	72.88	27.12	59.23	40.77
	M	31.84	68.16	27.08	72.92	27.12	72.88	43.47	56.53
	Classification Rate	68.23		66.22		72.88		57.88	
Gattal et al. (2020)	F	93.42	6.58	42.42	57.58	46.67	53.33	55.60	44.40
	M	12.77	87.23	18.79	81.21	19.26	80.74	14.80	85.20
	Classification Rate	90.32		61.82		63.70		70.40	
Kaljahi et al. (2019)	F	77.78	22.22	61.76	38.24	72.41	27.59	74.29	25.71
	M	40.35	59.65	37.50	62.50	64.00	36.00	52.00	48.00
	Classification Rate	68.72		62.13		54.20		61.14	

**Fig. 17.** Analyzing the results of the proposed and existing methods on different datasets for classification at line level.**Fig.18.** Failure cases of the proposed word segmentation step especially for Arabic and Farsi scripts.

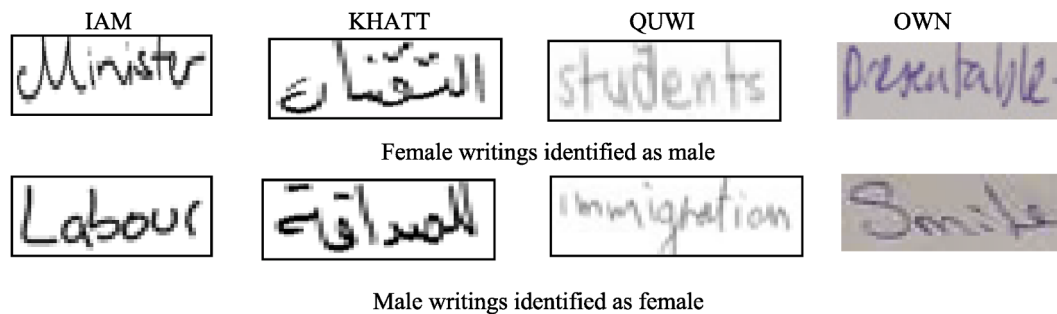


Fig.19. Misclassification of the proposed method for gender classification at word level.

the performance of gender identification. These are topics for future research.

CRedit authorship contribution statement

Shivakumara Palaiahnakote: Supervision, Writing – original draft. **Maryam Asadzadeh Kaljahi:** Formal analysis, Methodology. **Swati Kanchan:** Formal analysis, Methodology. **Umapada Pal:** Reviewing and editing. **Daniel Lopresti:** Reviewing and editing. **Tong Lu:** Validation, Visualization.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability

Data will be made available on request.

References

- Ahmed, M., Rasool, A. G., Afzal, H., & Siddiqi, I. (2017). Improving handwriting-based gender identification using ensemble classifiers. *Expert Systems with Applications*, 85, 158–168.
- Akbari, Y., Nouri, K., Sadri, J., Djeddi, C., & Siddiqi, I. (2017). Wavelet-based gender detection on off-line handwritten documents using probabilistic finite state automata. *Image and Vision Computing*, 59, 17–30.
- Arora, R., & Suman. (2012). Comparative analysis of classification algorithms on different datasets using WEKA. *International Journal of Computer Applications*, 54, 21–25.
- Banumathi, K. I. and Chandra, A. P. J (2016). "Line and word segmentation of Kannada handwritten text documents using projection profiles technique", In Proc. ICECCOT, pp 196-201.
- Bi, B., Suen, C. Y., Nobile, N., & Tan, J. (2018). A multi-feature extraction approach for gender identification of handwriting based on kernel mutual information. *Pattern Recognition Letters*.
- Bouadjene, N., Nemmour, H and Chibani, Y (2015). "Age, gender and handedness prediction from handwriting using gradient features", In Proc. ICDAR, pp 1116-1120.
- Bouadjene, N., Nemmour, H., & Chibani, Y. (2016). Robust soft-biometric prediction from off-line handwriting analysis. *Applied Soft Computing*, 46, 980–990.
- Ding, B., & Wen, G. (2019). Combination of global and local filters for robust SAR target recognition under various extended operating conditions. *Information Sciences*, 476, 48–63.
- Gattal, A., Djeddi, C., Siddiqi, I., & Chibani, Y. (2018). Gender classification from offline multi-script handwriting images using oriented basis image features (oBIFs). *Expert Systems with Applications*, 99, 155–167.
- Gattal, A., Djeddi, C., Bensefia, A and Ennaji, A (2020). "Handwriting based gender classification using COLD and Hinge features", In Proc. ICISP, pp 233-242.
- Kaljahi, M. A., Varshini, P. V. V., Shivakumara, P., Pal, U., Lu, T and Guru, D. S (2019). "Word-wise handwriting-based gender identification using Multi-Gabor response fusion", In Proc. DAR, pp 119-132.
- Khare, V., Shivakumara, P., Navya, N. J., Shwetha, G. C., Guru, D. S., Pal, U and Lu, T (2018). "Weighted gradient features for handwritten line segmentation", In Proc. ICPR, pp 3651-3656.
- Kingma, P. D and Bai, J. L (2015). "Adam: A method for stochastic optimization", In Proc. ICLR, pp 1-15.
- Lamsaf, A., Aitekerroum, M., Boulaknadel, S and Fakhri, Y (2019). "Text line and word extraction of Arabic handwriting documents", In Proc. SCA, pp 492-503.
- Liwicki, M., Schlapbach, A., & Bunke, H. (2011). Automatic gender detection using on-line and off-line information. *Pattern Analysis and Applications*, 14, 87–92.
- Louloudis, G., Gatos, G., Pratikakis, I., & Halatsis, C. (2009). Text line and word segmentation of handwritten documents. *Pattern Recognition*, 42, 3169–3183.
- Ma, F., Zhu, X., Wang, C., Liu, H., & Jing, X. Y. (2019). Multi-orientation and multi-scale features discriminant learning for palm print recognition. *Neurocomputing*, 348, 169–178.
- Maji, P., Chakraborty, S., Samanta, S., Chatterjee, S., Kausar, N and Dey, N (2015). "Effect of Euler number as a feature in gender recognition system from offline handwritten signature using neural networks", In Proc. INDIACOM, pp 1869-1873.
- Manyala, A., Cholakal, H., Anand, V., Kanhangad, V., & Rajan, D. (2018). CNN-based gender identification in near-infrared periocular images. *Pattern Analysis and Applications*.
- McAllister, P., Zheng, H., Bond, R and Moorhead, A (2016). "Towards Personalized Training of Machine Learning Algorithms for Food Image Classification Using a Smartphone Camera. In Proc. ICUCAL, pp 178-190.
- Mirza, A., Moetesum, M., Siddiqi, I and Djeddi, C (2016). "Gender classification from offline handwriting images using textural features", In Proc. ICFHR, pp 395-398.
- Moetesum, M., Siddiqi, I., Djeddi, C., Hannad, Y and Maadeed, S. A (2018). "Data driven feature extraction for gender classification using multi-script handwritten text", In Proc. ICFHR, pp 564-569.
- Nanni, L., Chidoni, S., & Brahnam, S. (2018). Ensemble of convolutional neural networks for bioimage classification. *Applied Computing and Informatics*.
- Narayan, S (1997). "The Generalized Sigmoid Activation Function: Competitive Supervised Learning", Information Sciences, pp 69-82.
- Nasr, G. E., Badr, E. A and Joun, C (2002). "Cross Entropy Error Function in Neural Networks: Forecasting Gasoline Demand", In Proc. FLAIRS, pp 381-384.
- Navya, B. J., Shwetha, G. C., Shivakumara, P., Roy, S., Guru, D. S., Pal, U and Lu, T (2018a). "Multi-gradient directional features for gender identification", In Proc. ICPR, pp 3657-3662.
- Navya, B. J., Shivakumara, P., Shwetha, G. C., Roy, S., Guru, D. S., Pal, U and Lu, T (2018b). "Adaptive multi-gradient kernels for handwriting based gender identification", In Proc. ICFHR, pp 392-397.
- Osman, Y (2013). "Segmentation algorithm for Arabic handwritten text based on contour analysis, In Proc. ICEEE, pp 447-452.
- Siddiqi, I., Djeddi, C., & Meslati, L. S. (2015). Automatic analysis of handwriting for gender identification. *Pattern Analysis and Applications*, 18, 887–899.
- Sokic, E., Salihbegovic, A and Djokic, M. A (2012). "Analysis of off-line handwritten text samples of different gender using shape descriptors", In Proc. BIHTEL, pp 1-6.
- Tan, J., Bi, N., Suen, C. Y and Nobile, N (2016). "Multi-feature selection of handwriting for gender identification using mutual information", In Proc. ICFHR, pp 578-583.
- Topaloglu, M., & Ekmekci, S. (2017). Gender detection and identifying one's handwriting with handwriting analysis. *Expert Systems with Applications*, 236–243.
- Wshah, S., Shi, Z and Govindaraju, V (2009). "Segmentation of Arabic handwriting based on both contour and skeleton segmentation", In Proc. ICDAR, pp 793-797.