

See discussions, stats, and author profiles for this publication at: <https://www.researchgate.net/publication/264044873>

The state of the art in handwriting synthesis

Conference Paper · December 2013

CITATIONS

7

READS

2,680

1 author:



Randa Elanwar

Electronics Research Institute

19 PUBLICATIONS 126 CITATIONS

SEE PROFILE

Some of the authors of this publication are also working on these related projects:



Physical layout analysis of scanned Arabic documents using machine learning techniques [View project](#)

The state of the art in handwriting synthesis

Randa I. Elanwar
Computers and Systems Dept.
Electronics Research Institute
Cairo, Giza, EGYPT
eng_r_i_elanwar@eri.sci.eg

Abstract: - Cursive handwriting is a complex graphic realization of natural human communication. Its production and recognition involve a large number of highly cognitive functions including vision, motor control, and natural language understanding. Handwriting synthesis has many important applications to facilitate user's work and personalize the communication on pen-based devices. The problem of handwriting synthesis is not new and a number of studies have been published in the literature. Some approaches, Movement-simulation techniques, make a real attempt at modeling the process of handwriting production. Other approaches, Shape-simulation techniques, which usually record the glyphs directly and reuse or sample the glyphs when synthesis. Different challenges are holding back the progress of such type of research. In this paper we present a literature review about the recent trends in handwriting synthesis, highlighting the different generation processes and pointing out the challenges facing the researchers. Finally we are giving a conclusion about the scientific research collections presented, and summarizing our opinions to help move future work up to maturity.

Key-Words: - cursive handwriting, datasets, generation, glyph-based, motor-models, synthesis

1 Introduction

Handwriting synthesis, i.e., converting ASCII text into the user's personal handwriting, is an important yet much less explored problem. Handwriting synthesis can be helpful to forensic examiners, the disabled, and the researchers working on handwriting recognition systems. Plentiful applications of handwriting synthesis can easily find their way in biometric security systems, information retrieval (web search) and natural language understanding. Furthermore, synthesized handwriting can also contribute to the personalization of one's computing devices like Tablet PCs and other pen-based interfaces.

Several methods have been reported in the literature for synthetic text generation. Though, these works address the problem of generating synthetic traits differently. As explained by [14] not all of them consider the term synthetic in the same way. In particular, three different strategies for producing synthetic samples can be found in the current literature:

i. Duplicated samples: In this case the algorithm starts from one or more real samples of a given person and, through different transformations, produces different synthetic (or duplicated) samples corresponding to the same person. This type of algorithm is useful to increase the amount of already

acquired handwriting data but not to generate completely new datasets. Therefore, its utility for performance evaluation and vulnerability assessment in handwriting is very limited. On the other hand, this class of methods can be helpful to synthetically augment the size of the enrollment set of signatures in identification and verification systems, a critical parameter in signature biometrics. The great majority of existing approaches for synthetic signature generation are based on this type of strategy.

ii. Combination of different real samples: This is the approach followed by most speech and handwriting synthesizers. This type of algorithms start from a pool of real n-grams (isolated letters, or combination of two or more letters) and using some type of concatenation procedure combine them to form the synthetic samples ([2] and [25]). Again, these techniques present the drawback of needing real samples to generate the synthetic handwriting trait and therefore their utility for performance evaluation and vulnerability assessment in handwriting is also very limited. As in the previous case, this perspective for the generation of synthetic data is useful to produce multiple handwriting samples of a given real user, but not to generate synthetic individuals.

iii. Synthetic-individuals: In this case, some kind of a priori knowledge about a certain handwriting trait is used to create a model that characterizes that handwriting trait for a population of subjects. New synthetic individuals can then be generated sampling the constructed model. In a subsequent stage of the algorithm, multiple samples of the synthetic users can be generated by any of the procedures for creating duplicated samples. Regarding performance evaluation and vulnerability assessment in handwriting this approach has the clear advantage over the two previously presented, of not needing any real handwriting samples to generate completely synthetic databases. This way, these algorithms constitute a very effective tool to overcome the usual shortage of handwriting data without undertaking highly resource-consuming collection/acquisition campaigns.

The existing methods for handwriting synthesis can be roughly divided into two categories [26]. The first one is based on the handwriting reconstruction process ([30] and [38]), in which the handwriting trajectory is analyzed and modeled by velocity or force functions, known as *movement simulation techniques*. Though physically plausible, these methods may not be convenient for synthesizing non-cursive handwriting since they mainly focus on the representation and analysis of real handwriting signals rather than handwriting synthesis. Extensive studies on motor models can give some insight but cannot directly provide a solution to the synthesis of a novel handwriting style. The second category involves glyph-based methods ([16] and [47]), also called *shape-simulation techniques*, which usually record the glyphs directly and reuse or sample the glyphs when synthesis. These methods require intensive user involvement in the sample collection process and cannot produce various handwriting styles, e.g., from handprint style to fully cursive style, in a natural way. Another issue of this approach is that it tries to learn a separate model for the connection of each pair of letters, and this is impractical since the size of the training set is always limited.

The synthesis process is not straightforward. Whether the motivation beyond handwriting synthesis is for creation or analysis, there is a number of common challenges that face researchers. These challenges are: i) the seed data samples from which glyph features are extracted or character models are designed, it should encounter enough variability for generalization, ii) the model

design and how to select and learn style-dependant parameters, iii) how to generate (deform) a synthesized sample that resembles the training set, iv) how to find out the best glyph model from synthesized samples to compose a word, v) how to concatenate the glyphs without introducing discontinuities to obtain cursive handwriting, vi) how to make the handwriting smoother to insure natural looking and vii) finally, how to evaluate the quality of the synthesized word.

The research in handwriting synthesis is moving forward but yet most of these issues are still open. In this paper, we present a comprehensive survey about handwriting generation in different fields of research and applications. Distinguishable sections illustrate the historical review of this research field, the recent trends, the existing challenges and research efforts done to overcome these difficulties. A final section stands on the observations and conclusions drawn from the scientific research collections presented, and summarizes our opinions to help move future work up to maturity.

2 Synthesis in biometric security

With the increasing importance that biometric security systems are acquiring in today's society and their introduction in many daily applications, a growing interest is arising for the generation of synthetic biometric traits such as voice, fingerprints, iris, handwriting, or signature. The generation of these synthetic samples is of interest, among other applications, for performance evaluation and vulnerability assessment of biometric systems. More specifically, synthetically generated biometric databases: i) facilitate the performance evaluation of recognition systems instead of the costly and time-consuming real biometric databases, and ii) provide a tool with which to evaluate the vulnerability of biometric systems to attacks carried out with synthetically generated traits [14].

Very limited work has been found in literature about handwriting synthesis for the purpose of biometric security. Of which is considered remarkable, the effort done by Galbally et. al [14] in 2009. They have proposed a method to generate synthetic online signatures. A first step, carried out in the frequency domain, in which the synthetic Discrete Fourier Transform (DFT) of the trajectory signals x and y is generated using a parametrical model, obtained by spectral analysis of a development set of real signatures. In the second

stage the resulting trajectory signals are used to place the penups of the pressure function. Finally, in the last stage, all the three signals are refined and processed in the time domain in order to give the synthetic signatures a more realistic appearance (smoothing, translation, rotation and scaling transformations are applied). Once a synthetic signature (defined by its x , y , and p functions) has been created, multiple samples of that master signature are generated (synthetic databases) applying the following deformations: Horizontal and vertical affine scaling of all three signals,

Duration expansion or contraction, and Noise addition (Smoothed white noise). The block diagram of the system is given in fig. 1. The performance of the synthetic signatures has been also evaluated using a Hidden Markov Model (HMM) based signature recognition system. Experiments have been conducted to see if the behavior of the synthetic signatures is comparable to those of the real samples, and thus can be used in the evaluation of signature verification systems.

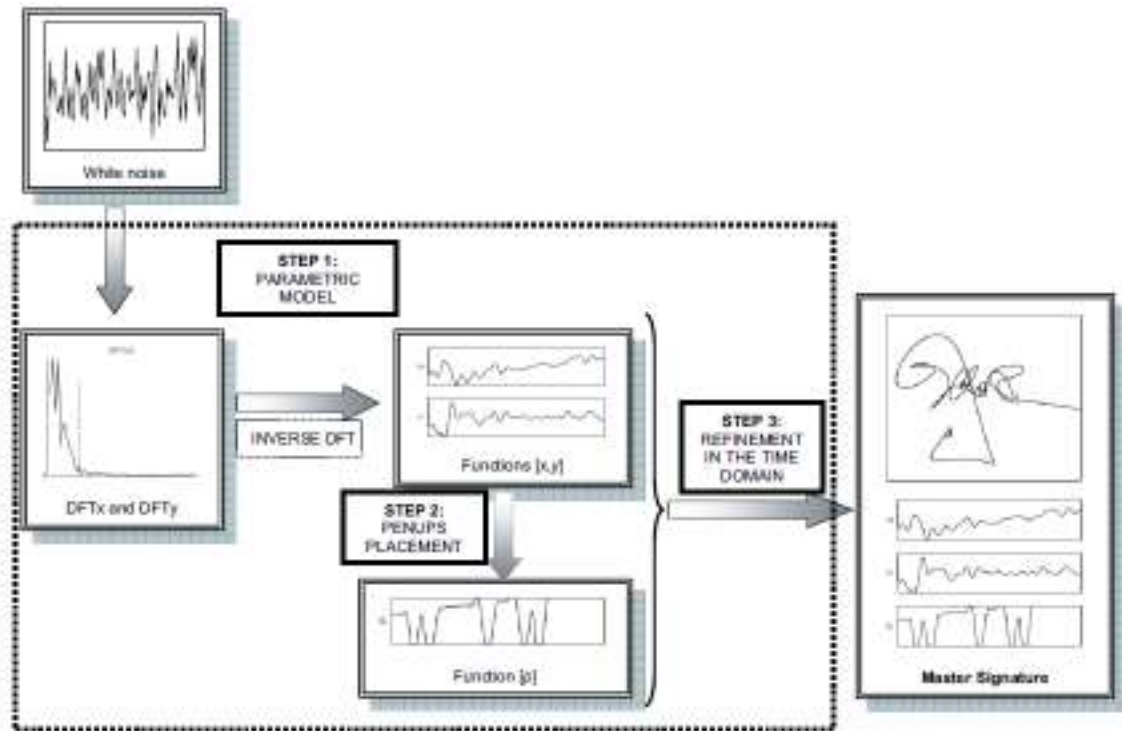


Fig. 1: Galbally et. al's system [14] for synthesizing signatures.

3 Synthesis in Information retrieval

Another important application for handwriting synthesis is information retrieval. Handwriting synthesis makes it possible to perform text searches on handwritten word image databases when no ground-truth data is available. This is still an open issue in information retrieval research field. The handwritten string is treated as a pictographic pattern without an attempt to understand it. This is a more natural way to handle the handwritten text. In this approach the query string is compared to database strings using an appropriate distance function. Unlike the conventional search techniques that work by comparing the ASCII codes of the ground-truth data and the user key words. This gives the user more expressive power; as the user can

extend his search to include also non-ASCII symbols (e.g., Latin), drawings and graphs and equations without being restricted by the key words language.

From the few publications found, that of the effort done by Rodriguez-Serrano et al in 2011 [35]. They have proposed a method to perform text searches on handwritten word image databases when no ground-truth data is available to learn models or select example queries. The approach proceeds by synthesizing multiple images of the query string using different computer fonts. The synthesized images undergo normalization and feature extraction operations. A word model is trained which makes use both (i) of the "synthesized" features, and (ii) of the feature distribution (vocabulary) obtained in the training

phase. The domain mismatch between queries (synthetic) and database images (handwritten) leads to poor accuracy. To solve this problem, they represent the queries with robust features and use a model that explicitly accounts for the domain mismatch. They use a semi-continuous hidden Markov model (SC-HMM). The key point of this model is that a part of the parameters can be estimated from handwritten data in an unsupervised way, while the remaining word-dependent parameters are directly learned from the synthetic examples. While the model is trained using synthetic images, its generative process produces samples according to the distribution of handwritten features. Finally, a font selection method (which is unsupervised, efficient, and keyword-independent) leverages the font contributions to best represent handwritten data, and yields significant improvements in accuracy. Experiments demonstrate that the proposed method is an effective way to perform queries without using any human annotated example in any part of the process.

4 Synthesis in handwriting recognition

Handwriting recognition has been an intensive research field for more than 40 years [3]. The most successful applications include isolated character recognition, automatic address processing, and bank check reading. However, for the problem of general unconstrained word and sentence recognition, where no lexicon is considered, the recognition rates are still rather low. Besides the recognition algorithm and the types of features used, the amount and quality of training data also has a great impact on the recognition performance of handwriting recognition systems. As a rule of thumb says, the classifier trained on the most data wins. This has been experimentally justified in many works before ([4], [37], [45]). The most straightforward way to expand the training set would be to collect additional natural, i.e. human written, samples. But collecting human written texts is a rather error prone, expensive and labor- and time- consuming process. Alternatively, the training set can be expanded by synthetic texts.

Several methods have been reported in the literature for synthetic text generation. Most of them generate the synthetic texts from existing natural ones. Some approaches use human written character tuples to build up synthetic texts [16]. In other

approaches, synthetic texts are generated by applying random perturbations on human written characters [4], words, or whole lines of text [43]. This has proved to be a very efficient way to increase the variability of the generated set of synthetic texts, yielding improvements of the recognition rate when used as additional training data to the handwriting recognition system ([4] and [43]).

However, generating synthetic texts does not necessarily require to use human written texts as a basis. In [24], Lee et. al. (1998) synthesize Korean characters using templates of characters, and a handwriting generation model. The templates consist of strokes of predefined writing order. After the geometrical perturbation of a template, beta curvilinear velocity and pen-lifting profiles are generated for the (overlapping) strokes. Finally, the character is drawn using the generated velocity and pen-lifting profiles. Although the generated characters look natural, they were not used to train a recognizer.

Similarly in [44], Varga et. al (2005) present an approach for the generation of English handwritten text lines by manually build templates (or prototypes) of handwritten characters using Bezier splines. As a first step to generate a text line image corresponding to a given ASCII transcription, the appropriate series of character templates are perturbed (scaling, shifting, and changing the slant) and concatenated. The resulting static image is then decomposed into strokes of straight segments and circular arcs. These strokes are then randomly expanded and overlapped in time. For each stroke, a delta-lognormal curvilinear velocity profile is generated [31], and the skeleton image of the text line is drawn, followed by grayscale thickening operations to make the generated script look more natural. The advantage of using a handwriting generation model is that the variations resulting from perturbing its parameters may better reflect psychophysical phenomena of human handwriting than applying geometric ad-hoc distortions on a static image. The text lines generated by the proposed method are used to train an HMM-based off-line handwritten text line recognizer to examine and evaluate how the synthesized training sets compare to the natural ones, in terms of the recognition rate and the recognizer capacity, i.e. the number of free parameters to be estimated.

Helmets et al. (2003) [17] has introduced a number of methods to generate synthetic handwritten text that can be used to train handwriting recognition systems. The basic idea is

to use image templates of single characters and n-tuples of characters of a certain writer either extracted from dataset of isolated characters or a dataset of cursive words or a dataset of n-tuples characters, and concatenate them to generate synthetic handwritten text of his own handwriting. All templates are produced by human writers. The methods proposed are experimentally evaluated in the context of an HMM based sentence recognizer that was developed earlier. The training data in experiments is a combination between natural data in addition to synthetic data varying according to the synthesis method. The recognition performance of the system based on synthetic training data of n-tuples characters slightly outperformed the system trained with natural data.

Chowdhury et al (2009) [9] present a complete sequential schema for a user-interactive personal

cursive handwriting reproduction system. As shown in fig. 2, during training, a handwritten sample from the user is processed, analyzed and recorded, and referred later for reproduction of the user's writing. Glyphs are efficiently segmented using user-interactive approach and fitted (represented) by individual cubic Bezier curves between a set of control points. For reconstruction, letter or bigrams glyphs are determined. Efficient realization of relative orientation of letters, x-alignment and y-alignment, runtime construction of inter-letter transition and maintenance of different writing styles of the same letter in different parts of the text is derived from the training dataset values distribution.

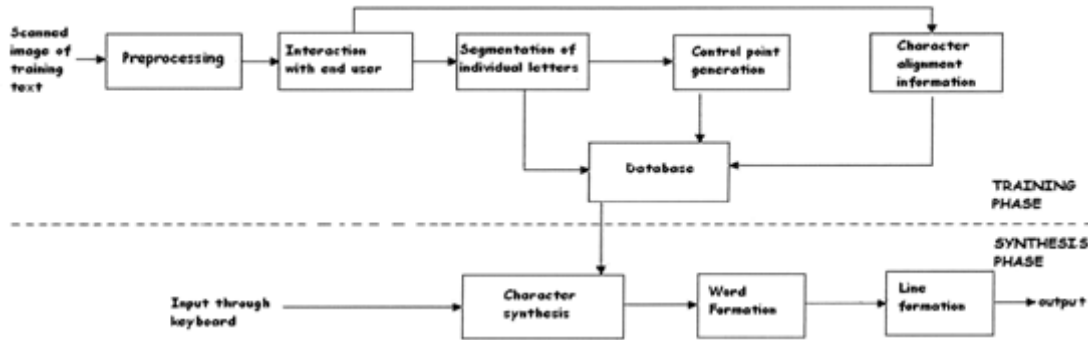


Fig. 2: Chowdhury's English cursive handwriting reproduction system

Elarian et al. (2011) [12] propose an approach to synthesize Arabic handwriting as shown in fig. 3. They segment word images into labeled characters and then use these in synthesizing arbitrary words. In Arabic some letters of the alphabet are connected to their within-word successors by a horizontal connection line called Kashidah. Kashidahs are easy to locate and recognize. They train the system by providing it with labeled segmented characters. The process of segmenting characters is semi-automated on the Kashidah position. Features from the connection lines (Kashidahs) cuts (i.e. to the right and to the left of the cut) are extracted and stored. Each letter sample is associated with a feature vector containing the features for whatever RK (Right Kashidah) or LK (Left Kashidah) is present. In the synthesis process, Kashidahs of the samples of the required word are compared. Starting from the first character, a chain of matching is conducted to find fitting Kashidahs according to the minimum distance measure. The character sample with RK matching best the LK of the previous character is

chosen. The chosen images are then aligned into words and sentences.

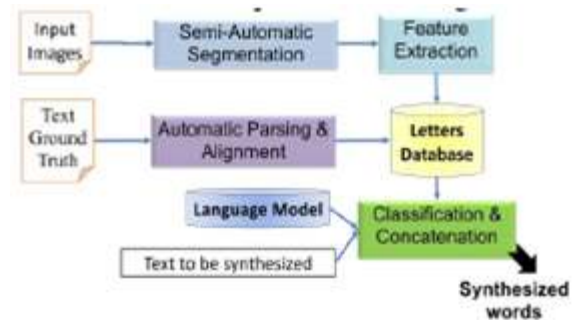


Fig. 3: Block diagram of Elarian's system to synthesize Arabic handwriting

Dinges et al (2011) [10] also propose a template based approach for generating synthesized handwritten Arabic characters. It starts by reconstructing a character trajectory then generating a polygonal representation for each training sample for every character class. Then for each class an Active Shape Model (ASM) is used to unify all

polygonal samples in one compact representation. Active Shape Models (ASMs) and novel Active Shape Structural Models (ASSMs) are used for generating variances of simple drawings and signatures. An ASM uses the distribution of some significant points to store the most important information of many shapes of a class in one single model. Accordingly, any desired number of synthesized characters, can be produced, as a result of simple linear combination between the Eigen values and Eigenvectors of the ASM. Ultimately, the contour of synthesized character is smoothed using piecewise cubic hermit interpolation. Finally, image representations (which should resemble the input images) are generated by a drawing algorithm. They claim that by combining multiple synthesized characters, the system is capable to produce synthesized Arabic words. They haven't neither achieved such goal nor have evaluated the system performance for isolated character synthesis.

More efforts can be found in literature basically for non-cursive letter glyphs ([1], [6], [27], [41], [46], and [50]). The deformation and cursiveness needed for word generation make synthesis more complicated task, since transitions from previous as well as to the next character are involved. Thus, we are only over viewing the cursive handwriting generation in this paper.

5 Synthesis in pen based computing

Pen-based interfaces are now a hotspot in Human-Machine Interface (HCI) research and the flourish of pen-based devices such as Tablet PCs brings a great demand for various cursive handwriting computing techniques. When writing a note on a Tablet PC, if the computer can automatically correct some written errors and generate some predefined handwriting strokes, communication through a pen-based interface would be more effective and intelligent. Furthermore, handwriting is preferable to typed text in some cases because it adds a personal touch. All these applications bring an urgent requirement for handwriting synthesis techniques.

This same idea can be applied for another important goal which is historical documents repairing. Invaluable historical documents may wear out by time and may cost a lot for expert consultancy and repair. Handwriting synthesis can help by learning the personal handwriting style and generating character glyphs to repair the damaged writing.

The first effort recorded in this field has been the work done in 1996 by Guyon [16]. She has introduced a straightforward approach to synthesize handwritten words. The system collects handwritten glyphs of single characters and letter groups that most frequently appear in English text, such as “tion” and “ing”. When synthesizing a word, the system splits the word into letter groups or characters. For example, “believe” may be partitioned into “be”, “li”, and “eve”. Then the corresponding glyphs are placed side by side without additional effort to connect them into a fluent handwriting. This method does not handle glyph variation (although a global transform is tried). As a result, the synthesized handwriting has a regular appearance, unexpected pen lifts exist and possible connections exist within each glyph only. In addition, the system requires users to write more than a thousand letter groups in order to provide complete samples, which is tedious and impractical.

From the outstanding later efforts, the work done by Jue Wang et al. in 2002 [47] and 2005 [48]. They perform a mathematical analysis of cursive handwriting style. A series of models are proposed to capture the characteristics of different (unique) writing styles. A tri-unit cursive handwriting model is proposed, which enables us to extract letter strokes and ligature strokes from cursive handwriting. Instantiations of letters are generated from the trained generative models, and some structural noise and transforms are added to increase the variability of the synthetic results. According to the generative model, the characteristics of a letter are determined by the locations of its control points. So the problem of learning the writing style of a letter is converted to learning the specific distributions of these control points (fig. 4). Each (x,y) component is viewed as an individual 1-D signal from which a series of multi-scale control points can be extracted from handwritten strokes using 1-D Gabor filters. B-splines are used to connect the control points to generate new samples of the letter. The trajectory of each letter can be divided into three parts in sequence: the head, the body and the tail units. The head unit of each letter is connected with and influenced by the tail unit of the previous letter and the tail unit is connected with and influences the head unit of the next letter. For a synthetic letter, its main body is generated from the letter style model, but its head part and tail part deserve some deformations to make a smooth and natural connection with neighbors. A concatenating energy and an energy minimization criterion are defined to guide the deformation process. Each

letter in a cursive handwriting is deformed in this way and a synthetic, style-preserving handwriting word finally formed. The synthesis frame work is given in fig. 5.



Fig. 4: (a) Control points extracted from the original strait. (b) Reconstructed strait using control points and B-splines

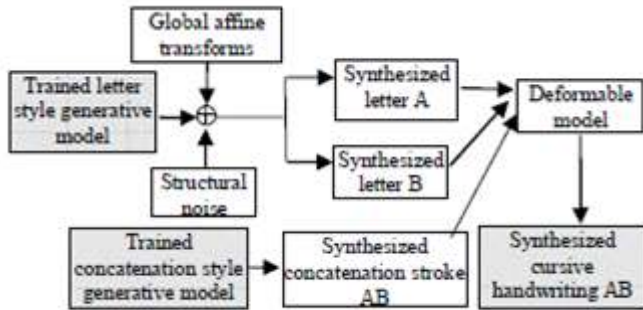


Fig. 5: Wang's cursive handwriting synthesis frame work

Another work done by Z. Lin and L. Wan in 2007 [26] to synthesize English handwriting in the user's writing style. They ask the user to input all individual characters three times, several special pairs of letters to extract features such as character glyph, size, slant, and pressure, and several multi-letter words to extract features such as special connection style, letter spacing, and cursiveness (fig. 6). Features are extracted as the user's writing style. After extracting the writing style, the system synthesizes cursive handwriting hierarchically.

The system assumes that only lowercase letters may be connected to each other. Lowercase letters are separated into uni-stroke letters and multi-stroke letters. The multistroke letters include "f", "i", "j", "t", "x", and "z", which are probably written in multiple strokes. Multi-stroke letters have more than one way of connecting to other lowercase letters.

For an input ASCII text, the glyphs of characters are first generated based on the features of handwriting style. The first letter glyph is chosen

based on the knowledge of its connection state, i.e., whether or not it is connected to its subsequent letters. For the remaining characters, the three samples are randomly selected. Then each glyph is geometrically perturbed. For stroke pieces delimited by high curvature points, we sequentially apply local random rotation and random scaling to them. The deformation is at a small scale so that the perturbed glyph looks similar to, but still different from, the original one.

To compose the glyph of a word, letter glyphs are aligned on the baseline with appropriate horizontal distance between neighboring glyphs and vertical offsets from the baseline (ascenders and descenders). Next, the adjacent letters are connected to each other using high-order polynomial interpolation. The heads or tails of the glyphs may be trimmed in order to avoid severe overlap and to ease connection. Then the pressure (stroke width variation) is assigned to the ligature, and words are rendered one by one to form a line. The lines are further stacked into paragraphs. The system flow chart is given in fig. 7.

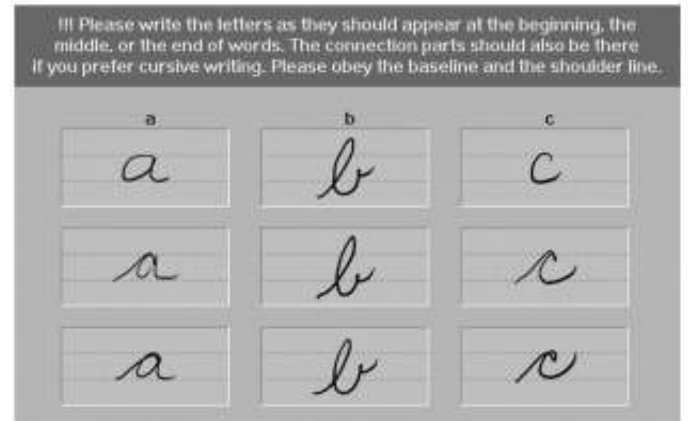


Fig. 6: The user interface (UI) to collect user handwriting samples.

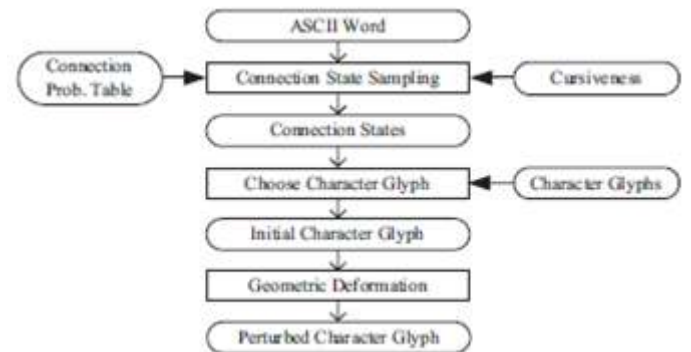


Fig. 7: Lin's system flow chart

More efforts can be found in literature for isolated glyph generation for online handwriting as in ([5], [8], [19], [20], [21], and [42]).

6 Conclusion

6.1 Summary of efforts

The handwriting of different people differs in many aspects. These aspects actually define the handwriting style of a person. As suggested by handwriting analysis techniques in forensic inspection or character analysis, factors that are easily noticeable to ordinary people to distinguish different handwriting styles include: 1. the glyph and the size of characters; 2. the pressure distribution and the slant of handwriting; 3. The relative sizes of the middle, the upper, and the lower zones of letters; 4. the existence and the shape of lead-in, connecting, and ending parts; 5. the letter, the word, and the line spacings; 6. the embellishment; and 7. the simplified or neglected strokes [26].

The problem of handwriting synthesis is not new. However, few studies starting from the late sixties have been published in the literature analyzing and studying the human handwriting moments [49] and the human "motor code" for the production of cursive handwriting ([15], [18], [30], [31], [39]).

The straightforward approach has been usually to synthesize handwriting from collected handwritten glyphs. Each glyph is a handwriting sample image of one, two or three letters. When synthesizing a long word, several glyphs are simply juxtaposed in sequence and not usually even connected to generate fluent handwriting or are connected using simple curves or polynomials causing unnatural looking for the handwriting due to overlaps or discontinuities.

The recent emergence of pen computers with high resolution tablets has made available dynamic (temporal) information as well as created the need for robust online handwriting recognition algorithms. Considerable effort has been spent on online recognition, but there are no robust, low error rate recognition schemes available yet.

From here people began to solve inverse problem (analysis by synthesis) and synthesis evaluation by recognition. They don't actually produce glyphs but they share the same idea of synthesis as they try to understand and model/feature the motor activity of the hand for example [13].

Even this type of research has passed by stages, stage 1 when the tablets were not yet able to record the pressure information and stage 2 when pressure information can be captured and involved into the problem solution.

For the glyph model creation researchers are usually trying to find a reliable description of a character shape through the extraction of some descriptive features like size [26], slant [26], and pressure [26]. Control points (local maxima of a filter) ([9], [47], and [48]), points of maximum/minimum x,y coordinate, and crossing points [48]. Some researchers use some new features related to the within-word connection style: includes the head connection type and the tail connection type, letter spacing: defined as the distance between the central lines of neighboring letters, and cursiveness: a measure of how much the user prefers cursive writing. It is between 0 and 1, where 0 represents that the user prefers handprint writing while 1 denotes that the user likes completely cursive writing [26].

For natural looking of the handwritten script, there should be some variation introduced to the different samples used for the same character glyph according to its position within the word. Writer variation is achieved via geometric perturbation stages (random local rotation-scaling) (as in [26]) or energy minimization of the extracted pen trajectory ([47] and [48]).

One more method is followed to achieve natural handwriting looking, which is introducing pressure signal to the online handwriting (pen-ups) and similarly varying the ink thickness along the written stroke in offline handwriting ([16] and [44]). Introducing pressure according to [14] in signature synthesis is achieved differently. The authors claim the two main features defining the pressure function of a signature are: i) the number of penups (i.e., zero pressure segments of the signature) that occurred during the signing process, and ii) the placing of those penups. The distribution of the number of penups was extracted from the real signature database, and applied to the synthetic signatures according to their length (i.e., a longer signature presents a higher probability of having a large number of penups). They claim that most penups occur close to a singular point (maximum or minimum). Thus, the penups are located through the pressure function and some maximum points (between penups) are determined randomly. In a successive step all these singular points (penups and maxima) are joined using a cubic spline interpolation algorithm.

In [26], the authors introduce pressure and use normalization scaling algorithm as well so that the size variance among the characters is minimized, and on the other hand the scaling factor for each sample is also close to 1.

The last stage in a handwriting synthesis is usually an evaluation stage to assess the performance of the procedures used. Evaluation is most probably done through visual examination by an expert to determine consistency ([12], [26], [47], and [48]). Another trend for evaluation is through using a HMM based handwritten text recognizer to evaluate how the synthesized training sets affect the performance, in terms of the recognition rate and the recognizer capacity ([14], [17], and [44]).

6.2 Challenges

Several methods have been reported in the literature for synthetic text generation. Most of them generate the synthetic texts from existing natural ones. In some approaches, synthetic texts are generated by concatenating natural human written glyphs (characters or n-tuples) to compose words ([9], [12], [16], [17]). Or differently, synthetic texts are generated by applying random perturbations on human written characters ([4], [10], [14], [26], [43], [47], [48]). Other approaches do not necessarily require to use human written texts as a basis ([35] and [36]). Instead, characters are synthesized using templates of characters, and a handwriting generation model. The templates consist of either standard fonts ([35] and [36]) or strokes of predefined writing order ([24] and [44]). After the geometrical perturbation of a template, velocity and pressure profiles are generated for the (overlapping) strokes. Finally, the character is drawn using the generated profiles.

For the approaches depending on the existence of natural human handwriting samples, the main challenge is to collect a large enough dataset of natural human handwritings encountering a wide range of writer variability. In the same time this violates the main target beyond synthesis, which is, generating large amount of data independent from natural sources.

Another consequent challenge is the modeling as depicted in [13]: One of the major difficulties of the cursive word recognition descends from the great variability observed in different samples of script issued from the same writer over time or from different scriptors. So it is difficult to find a reliable description of a word able to represent all the admitted occurrences of the input shape.

This issue raises the question about the target of the synthesis process, shall the problem be that we duplicate the data base i.e., understand the writing style and enlarge the data with the same style or we mix styles and get a new synthetic style to add different varieties? In this case, the performance becomes limited by samples used for training since the shape models can only generate novel shapes within the variation of training samples. To produce more variant and natural handwriting, users are required to give more handwriting samples, and sometimes this is not practical.

As a solution to increase the variability of synthetic data added to a recognizer training set for example, the distortion should not be too weak. On the other hand, severely distorted handwriting may bias the recognition system used for synthesis evaluation toward unnatural handwriting styles, which can lead to a deterioration of the recognition accuracy. So it is very important to find an optimal distortion parameter range and control the variability of data to ensure that the synthetic handwriting is natural enough.

Another important challenge, related to the natural looking of the synthesized data, that has appeared strongly is the handwriting cursiveness. As not all users tend to write all the words in cursive manner, on the contrary, the handwriting is most probably in a mixed style. Thus, during synthesis, the system has to determine which pair of adjacent letters in a word is connected. It would be ideal if we compute the connection probability from handwritten samples. However, it is impractical in real practice due to the large amount of letter pairs and again we may not have enough natural handwriting samples. Because writing all the pairs once is laborious, and writing each letter pair only once cannot provide an accurate estimate of the connection probability. Another area to be addressed is that the spline based system can only handle fluid cursive handwriting and must be coordinated with other techniques to deal with mixed-style handwriting, which may compose of abundant straight lines. Another issue is that the success of the system heavily depends on the accuracy of the segmentation of handwriting samples used for training open issues

Finally, how to give an objective evaluation on synthetic handwriting is still a challenging problem. Most of the reviewed work is either tending not to evaluate the quality of the synthesis process or to use human expertise rather than a strongly mathematical founded approaches. The only machine solution has been that of using HMM

recognizers to justify the improvement (if any) shown in recognition accuracy due to adding the synthetic handwriting to the training set of the recognizer.

6.3 Opinions

There are several ways to generate synthetic handwriting, e.g., the motor-based method and template (glyph) based synthetic handwriting generation. However, most of them are applied to isolated characters. Fewer research efforts have been directed to cursive handwriting synthesis. We have reviewed the latest and the most remarkable efforts done for handwriting generation in different scientific fields. We have summarized the approaches used and the challenges faced. Consequently, we have stood on (i) the motivations beyond this field of research and its importance to other scientific branches, (ii) the synthesis process stages and the system inputs and functioning, (iii) the techniques the researchers use to achieve their goal, and finally (iv) the obstacles they face whether or not they have succeeded to overcome.

From the presented efforts in this paper we have had some remarks that may help move future work up to maturity and improve the quality of the synthesized handwritten scripts. These remarks can be summarized as:

The main problem and challenge is the need of large amount of handwritten samples for training the glyph models in order to generalize the models and avoid biasing to a limited set of writers variety. The important issue that we have noticed that all researchers are either tending to use local datasets of their own rather than publically available handwriting datasets or using very limited part of a public dataset. This may be attributed to their tendency to use human expertise for evaluation which will be very laboring to handle in case of relatively large datasets. The only solution for this problem is to find out an objective and reliable automatic evaluation scheme in order to accelerate the process and allow using much bigger dataset.

Further, most researchers are tending to collect isolated glyph samples which leads to unnatural looking at synthesizing words. This problem will be almost solved if they use cursive words for training. Here, automation is also needed for training dataset preprocessing, specially word to character segmentation to help enlarging the dataset used and provide more variety and generalization.

In the generation process, the glyph model produces samples according to the distribution of

handwritten features extracted from the whole training dataset. Regardless the dataset size, such method for a model construction in case of large variances existing in some feature values may deform the model and lead to odd looking generated samples. Instead, writer clustering relative to their handwriting scheme and generating different models for the same character per cluster will help enhance the quality of the generated samples and restrict the distortion range. In other words, this leads to a combination between the motor-based methods and the glyph based methods for handwriting synthesis which will unquestionably improve the generation results getting the benefits of the two approaches together.

Another issue considering the modeling is the features used. It has been noticed that most researchers tend to use almost the same features or subset of them (geometric, polynomials, filter control points, etc.) which may be a good reason for limited performance. As a suggestion, in case of no other novel descriptive features could be thought of, feature fusion techniques may significantly help change the distribution in feature space. Consequently, modeling will be easier and better.

Finally, all the researchers tending to evaluate their synthesis process are using recognizers (especially HMM). They all tend to add the synthetic datasets to the training data and test their system by the recognition performance for a natural handwritten data set. As a matter of fact, there another way of thinking for using recognizers. Recognizers can be used to enhance synthesis not to evaluate. The training data set has to be natural handwritten data and the test dataset should be synthetic. In such case, better recognition performance will be targeted in order to tune the distortion parameter range and modify the synthesized glyph samples look. Reaching a proper recognition rate, may then qualify the synthetic dataset to be added as training data for another recognizer type for evaluation.

References:

- [1] Atanasio V., Allographic Biometrics and Behavior Synthesis, *TUGboat, Proceedings of EuroTEX 2003*, 2003, Volume 24, No. 3, pp. 328-333.
- [2] Ballard, L., Lopresti, D., and Monroe, F., Forgery quality and its implications for behavioral biometric security, *IEEE Trans. on Systems, Man, and Cybernetics*, 2007, 37, 1107-1118.

- [3] Bunke, H., Recognition of Cursive Roman Handwriting - Past, Present and Future, *Proc. 7th Int. Conf. on Document Analysis and Recognition*, 2003, pp. 448-461.
- [4] Cano, J., Perez-Cortes, J.C., Arlandis, J. & Llobet, R., Training Set Expansion in Handwritten Character Recognition, *Proc. 9th SSPR / 4th SPR*, 2002, pp. 548-556.
- [5] Chang W., Shin J., A statistical handwriting model for style-preserving and variable character synthesis, *IJDAR*, 2012, 15:1–19
- [6] Chaudhuri B. B. and Kundu A., Synthesis of individual handwriting in Bangla script, *Proc. of 1st International Conference on Frontiers in Handwriting Recognition, ICFHR 2008*, Montreal, Canada, 2008.
- [7] Cheng W. and Lopresti D., Parameter Calibration for Synthesizing Realistic-Looking Variability in Offline Handwriting, *Document Recognition and Retrieval XVIII, part of the IS&T-SPIE Electronic Imaging Symposium, San Jose, CA, USA, Proceedings. SPIE Proceedings 7874 SPIE 2011*, 2011, pp. 1-10, ISBN 978-0-8194-8411-6, 2011
- [8] Choi H., Cho S. J., and Kim J. H., Writer Dependent Online Handwriting Generation with Bayesian Network, *Proceedings of the 9th Int'l Workshop on Frontiers in Handwriting Recognition (IWFHR-9 2004)*, 2004.
- [9] Chowdhury S., Das S., Roy D., Sarkar U. and Chaudhuri B. B., A complete method of personal handwriting synthesis, *Advanced in Graphonomics: Proceedings of IGS 2009*, 2009, pp. 250-253
- [10] Dinges L., Elzobi M., Al-Hamadi A., and Al Aghbari Z., *Synthesizing Handwritten Arabic Text Using Active Shape Models*, R.S. Choraś (Ed.): Image Processing & Communications Challenges 3, AISC 102, 2011, pp. 401–408.
- [11] Djioua, M., O'Reilly, C., and Plamondon, R., An interactive trajectory synthesizer to study outlier patterns in handwriting recognition and signature verification, in *Proc. ICPR*, 2006.
- [12] Elarian, Y.S., Al-Muhsateb, H.A., Ghouti, L.M.: Arabic handwriting synthesis. In: *First International Workshop on Frontiers in Arabic Handwriting Recognition*, 2011, <http://hdl.handle.net/2003/27562>
- [13] Elbert B. L., *Analysis by synthesis in Handwriting*, Master thesis, MIT, 1997.
- [14] Galbally J., Fierrez J., Martinez-Diaz M., and Ortega-Garcia J., Synthetic Generation of Handwritten Signatures Based on Spectral Analysis, *Biometric technologies for human identification VI, Proceedings of SPIE*, 2009.
- [15] Guerfali W. and Plomondon R., The Delta LogNormal theory for the generation and modelling of handwriting recognition, In *Proceedings of ICDAR'95*, volume I, 1995, pp. 495-498.
- [16] Guyon I., Handwriting synthesis from handwritten glyphs, In *Proceedings of the Fifth International Workshop on Frontiers of Handwriting Recognition*, 1996, pp. 309-312.
- [17] Helmers M. and Bunke H., *Generation and Use of Synthetic Training Data in Cursive Handwriting Recognition*, F.J. Perales et al. (Eds.): IbPRIA 2003, LNCS 2652, 2003, pp. 336-345.
- [18] Hollerbach J.M., An oscillation theory of handwriting, *Biological cybernetics*, 1981, pp 139-156.
- [19] Jawahar C. V., Balasubramanian A., Synthesis of Online Handwriting in Indian Languages, *Proceedings of International Workshop on Frontiers in Handwriting Recognition*, La Baule, France, 2006.
- [20] Jawahar C.V., Balasubramanian A., Meshesha M., Namboodiri A., Retrieval of online handwriting by synthesis and matching, *Pattern Recognition*, 2009, 42:1445--1457
- [21] Kamada M., Enkhbat R. and Kawahito T., Synthesis of Letter Strings in Script Style Based on Minimum Principle, *Advanced Modeling and Optimization*, 2004, 6(1), pp. 37-56.
- [22] Kamel I., On Indexing Handwritten Text, *International Journal of Multimedia and Ubiquitous Engineering*, 2010, Vol. 5, No. 2, pp. 39-53.
- [23] Kim J. H. and Choi H.I., Co-articulation Modeling for Handwriting Synthesis, *Proceedings of the International Conference on Computing: Theory and Applications (ICCTA'07)*, 2007.
- [24] Lee, D.-H. & Cho, H.-G., A New Synthesizing Method for Handwriting Korean Scripts, *Int. Journal of Pattern Recognition and Artificial Intelligence*, 1998, 12(1), pp. 46-61
- [25] Lin, A. and Wang, L., Style-preserving english handwriting synthesis, *Pattern Recognition*, 2007, 40, 2097-2109.
- [26] Lina Z. and Wan L., Style-preserving English handwriting synthesis, *Pattern Recognition*, 2007, 40:2097 – 2109
- [27] Liu G., Jin L., Ding K., Yan H., A new approach for synthesis and recognition of large scale handwritten Chinese words, *12th*

- International Conference on Frontiers in Handwriting Recognition*, 2010, pp. 571-575
- [28] Munich, M. E. and Perona, P., Visual identification by signature tracking, *Trans. PAMI*, 2003, 25:200-217.
- [29] Oliveira, C., Kaestner, C. A., et al., Generation of signatures by deformations, in *Proc. BSDIA*, 1997, 283-298.
- [30] Plamondon, R., A Kinematic Theory of Rapid Human Movements. Part I: Movement Representation and Generation, *Biological Cybernetics*, 1995, 72, pp. 295-307.
- [31] Plamondon, R. & Guerfali, W., The Generation of Handwriting with Delta-lognormal Synergies, *Biological Cybernetics*, 1998, 78, pp. 119-132.
- [32] Plamondon R., Li X., Djiova M., Extraction of delta-lognormal parameters from handwriting strokes, *Front. Comput. Sci. China*, 2007, 1(1): 106–113
- [33] Rabasse, C., Guest, R. M., and Fairhurst, M. C., A method for the synthesis of dynamic biometric signature data, in *Proc. ICDAR*, 2007.
- [34] Richiardi, J., *Skilled and synthetic forgeries in signature verification: large-scale experiments*, tech. rep., Institute of Electrical Engineering, Swiss Federal Institute of Technology, Lausanne, 2008.
- [35] Rodriguez-Serrano J. A., and Perronnin F., Synthesizing queries for handwritten word image retrieval, *Pattern Recognition*, 2012, 45(9): 3270-3276.
- [36] Rodriguez-Serrano J. A., and Perronnin F., Handwritten word-image retrieval with synthesized typed queries, *10th International Conference on Document Analysis and Recognition*, 2009, pp. 351-355.
- [37] Rowley, H. A., Goyal, M. & Bennett, J., The Effect of Large Training Set Sizes on Online Japanese Kanji and English Cursive Recognizers, *Proc. 8th Int. Workshop on Frontiers in Handwriting Recognition*, 2002, pp. 36-40.
- [38] Schomaker L.R.B., *Simulation and recognition of handwriting movements*, Doctoral Dissertation/Ph.D. Thesis (NICI TR-91-03), Nijmegen University, The Netherlands, 1991.
- [39] Singer Y. and Tishby N., Dynamical encoding of cursive handwriting. In *IEEE conference of computer vision and pattern recognition*, 1993.
- [40] Gabriel Terejanu, Handwriting Synthesis for Human Interactive Proof in Web Services, CSE 717 Project – Progress Report.
- [41] Thomas A. O., and Govindaraju V., Generation and Performance Evaluation of Synthetic Handwritten CAPTCHAs, *Proc. of 1st International Conference on Frontiers in Handwriting Recognition, ICFHR 2008*, Montreal, Canada, 2008.
- [42] Uchida S., *Statistical Deformation Model for Handwritten Character Recognition, Recent Advances in Document Recognition and Understanding*, Dr. Minoru Mori (Ed.), 2011, ISBN: 978-953-307-320-0, InTech, Available from: <http://www.intechopen.com/books/recent-advances-in-document-statistical-deformation-model-for-handwritten-character-recognitionrecognition-andunderstanding/>
- [43] Varga, T. & Bunke, H., On-line Handwriting Recognition Using Synthetic Training Data Produced by means of a Geometrical Distortion Model, *Int. Journal of Pattern Recognition and Artificial Intelligence*, 18(7), 2004, pp. 1285-1302.
- [44] Varga T., Kilchhofer D. and Bunke H., Template-based Synthetic Handwriting Generation for the Training of Recognition Systems, *Advances in Graphonomics: Proceedings of IGS 2005*, 2005.
- [45] Velek, O. & Nakagawa, The Impact of Large Training Sets on the Recognition Rate of On-line Japanese Kanji Character Classifiers, *Proc. 5th Int. Workshop on Document Analysis Systems*, 2002, pp. 106-109.
- [46] Vincent N., Seropian A., Stamon G., Synthesis for handwriting analysis, *Pattern Recognition Letters*, 2005, 26:267–275.
- [47] Wang J., Wu C., Xu Y., Shum H. and Ji L., Learning-based Cursive Handwriting Synthesis, in *Proc. of IWFHR*, 2002, pp. 157-162.
- [48] Wang J., Wu C., Xu Y., Shum H., Combining shape and physical models for online cursive handwriting synthesis, *IJDAR*, 2005, 7: 219–227
- [49] Yasuhara M., Handwriting Analyzer and Analysis of human handwriting moments, *Japanese Psychological Research*, vol. 11, No. 3, 1969, pp. 103-109.
- [50] Zheng Y., and Doermann D., Handwriting Matching and Its Application to Handwriting Synthesis, *ICDAR*, 2005, pp. 861-865.
- [51] Zheng Y., *Handwriting Identification, Matching, And Indexing In Noisy Document Images*, PhD Thesis, University of Maryland, 2005.