

Article

Basic Test Framework for the Evaluation of Text Line Segmentation and Text Parameter Extraction

Darko Brodić ^{1,*}, Dragan R. Milivojević ² and Zoran Milivojević ³

¹ Technical Faculty Bor, V.J. 12, University of Belgrade, 19210 Bor, Serbia;

E-Mail: dbrodic@tf.bor.ac.rs

² Department of Informatics, Zeleni Bulevar 35, Mining and Metallurgy Institute, 19210 Bor, Serbia;

E-Mail: dragan.milivojevic@irmbor.co.rs

³ Technical College Niš, Aleksandra Medvedeva 20, 18000 Niš, Serbia;

E-Mail: zoran.milivojevic@vtsnis.edu.rs

* Author to whom correspondence should be addressed. E-mail: medijum@verat.net;

Tel.: +381-30-424-555; Fax: +381-30-421-078.

Received: 9 April 2010; in revised form: 13 May 2010 / Accepted: 14 May 2010 /

Published: 25 May 2010

Abstract: Text line segmentation is an essential stage in off-line optical character recognition (OCR) systems. It is a key because inaccurately segmented text lines will lead to OCR failure. Text line segmentation of handwritten documents is a complex and diverse problem, complicated by the nature of handwriting. Hence, text line segmentation is a leading challenge in handwritten document image processing. Due to inconsistencies in measurement and evaluation of text segmentation algorithm quality, some basic set of measurement methods is required. Currently, there is no commonly accepted one and all algorithm evaluation is custom oriented. In this paper, a basic test framework for the evaluation of text feature extraction algorithms is proposed. This test framework consists of a few experiments primarily linked to text line segmentation, skew rate and reference text line evaluation. Although they are mutually independent, the results obtained are strongly cross linked. In the end, its suitability for different types of letters and languages as well as its adaptability are its main advantages. Thus, the paper presents an efficient evaluation method for text analysis algorithms.

Keywords: OCR; document engineering; text line segmentation; text features; testing

1. Introduction

Printed text is defined by strong shape regularity. Its text lines have similar orientation and its skewness is also similar or equal, hence text orientation on same page is not variable. Descenders and ascenders from neighbor text lines are mostly disjoint and consequently, they do not interfere. Accordingly, text distances between lines are big enough to regularly split up text lines. Word in text lines are formed regularly, with similar distances and inter word spacing is decent.

Handwritten text is fully or partially cursive text. It tends to be multi-oriented and skewed. Text lines in handwritten documents are primarily curvilinear and close to each other. Descenders and ascenders from neighbor text lines are occasionally mixed up. Text distances between lines are close to each other, hence text lines run in to each other. Words in text lines are not formed regularly, so their distance is different. On the other hand, like printed text, handwritten text inter-word spacing is tolerable. Overall the appearance of skewed lines with different orientation and text lines close to each other make handwritten text less readable.

From the above, printed and handwritten text are characterized by their feature diversity. Hence, their text line segmentation as well as parameter extraction procedure can be quite dissimilar, although algorithms should fulfill these tasks for printed as well as for handwritten text.

Prior to text parameter extraction, text line segmentation should be done. It is an important step in document image processing. Although some text line detection techniques are successful in printed documents, processing of handwritten documents has remained a key problem in OCR [1,2]. Most text line segmentation methods are based on the assumptions that the distance between neighboring text lines is significant and that text lines are reasonably straight. However, these assumptions are not always valid for handwritten documents. Hence, text line segmentation is a leading challenge in document image analysis [3].

Upon completion of this process, the primary goal of OCR is the extraction of text parameters from optically scanned documents, so reference text line and skew rate identification is mandatory. Their validity is of major importance for any OCR process. There are various reasons for the appearance of multi-skewed lines in text, but two of them are the most common [1]: Firstly, some degree of misalignment of the document during the scanning process is unavoidable, but since all printed text lines in the scanned document are uniformly skewed, this way the reference text lines are almost parallel. Secondly, text lines in an original handwritten document are skewed differently due to specific individual handwriting habits, so handwriting text lines present different orientations, *i.e.*, they are multi-skewed. To enhance the ability of document analysis system, we need a robust algorithm for text line segmentation as well as for parameter extraction.

Many proposed algorithms have been evaluated by quite different test methods. In fact, these evaluation procedures are usually based on use of a custom text database as a test sample. Accordingly, testing result interpretation is quite dissimilar [4]. Hence, the establishment of the test framework for the evaluation of the document image processing algorithms is of great importance. This is precisely the task of this paper, and a basic method framework for the evaluation of the text line segmentation and text parameters extraction is proposed.

The paper is organized as follows: in Section 2 the test framework is presented. It is divided into two test groups. Each of them is completely described. Section 3 contains an examination and

evaluation of the test framework procedure using an example of the specific algorithm. A Gaussian isotropic kernel is used as the basic test algorithm. Results are analyzed, examined and discussed. Conclusions are given in Section 4.

2. Evaluation Test Framework

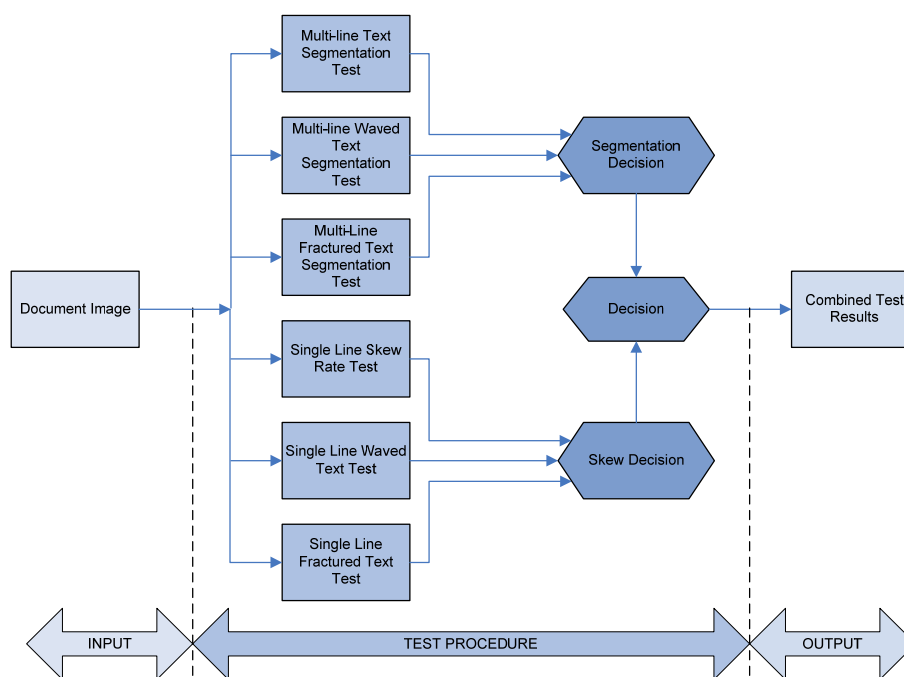
The evaluation test framework for the text parameter extraction algorithm consists of a few text experiments. They are divided into two distinct groups:

1. Text line segmentation experiments,
2. Reference line and skew rate experiments.

Text line segmentation experiments are related to the algorithm's ability to achieve segmentation of the text lines. Hence, these experiments are based on various multi-line sample texts. They incorporate the following tests:

- Multi-line text segmentation test,
- Multi-line waved text segmentation test,
- Multi-line fractured text segmentation test.

Figure 1. Integral framework test procedure.



In contrast, the reference line and skew rate tests evaluate the algorithm's competence for text line tracking. Therefore, they are based on single line sample texts as a reference. They include the following tests:

- Single line skew rate test,
- Single line waved text test,
- Single line fractured text test.

A schematic diagram of the integral framework test procedure is shown in Figure 1. According to everything presented above, decision making is required at the end of the test procedure. Firstly, decision making is mandatory for the multi-line text segmentation experiments process. As a result of this decision, the sub-set values of the algorithm parameters are obtained. These parameter values are used as an optimization starting point. Further, results from single line text experiments are evaluated. These results narrow the algorithm optimization choice by creating its own parameters value sub-set. Although the test experiments are quite diverse, their results are inter-related. Hence, the last decision includes a new parameter sub-set values taking into account the all previously obtained parameter sub-set values. This final result represents the optimized parameter values.

2.1. Document text image

At the beginning of the test process, an original image is used. Assume that original image is continual function $f(x, y)$. A document text image is obtained as a product of the original image scanning. Hence, the values of the coordinates (x, y) become discrete quantities. Now, the document text image is a digital text image represented by a matrix $\mathbf{D}$ with M rows, N columns, and intensity with L discrete levels of gray. L is the integer number from the set $\{0, \dots, 255\}$. Hence, the intensity of matrix $\mathbf{D}$ is represented as [5]:

$$D(r, c) = f(x, y) \quad (1)$$

where the origin of the function $f(x, y)$ is point $(x, y) = (0, 0)$, while the origin of the matrix $\mathbf{D}$ is $(r, c) = (1, 1)$. Hence, row $r \in \{1, \dots, M\}$ replaces $x \in \{0, \dots, M-1\}$ and column $c \in \{1, \dots, N\}$ replaces $y \in \{0, \dots, N-1\}$.

After applying intensity segmentation with binarization, the intensity function is converted into a binary intensity function given by:

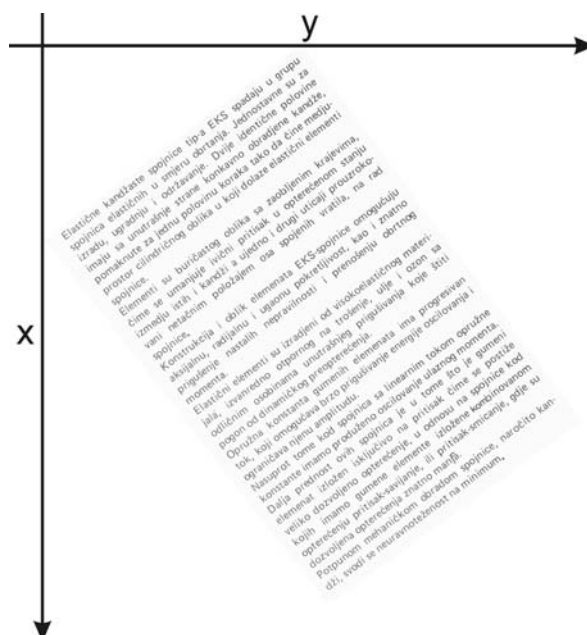
$$D_{bin}(r, c) = \begin{cases} 1 & \text{for } D(r, c) \geq D_{th} \\ 0 & \text{for } D(r, c) < D_{th} \end{cases} \quad (2)$$

where D_{th} is given by the Otsu algorithm [6]. It represents a threshold sensitivity decision value.

Now, extracted text lines are represented as a digitized document image by matrix $\mathbf{X}$ featuring M rows by N columns. Currently, the document text image is represented as a black and white image. It consists of the only black and white pixels. Each character or word consists of the only black pixels. Each pixel $\mathbf{X}_{ij}$, i.e., $\mathbf{X}(i, j)$ is represented by the number of coordinate pairs such as:

$$\mathbf{X}(i, j) \in \{0, 255\} \quad (3)$$

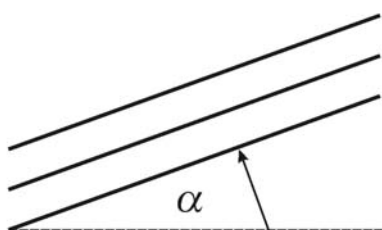
where $i = 1, \dots, M, j = 1, \dots, N$ of matrix $\mathbf{X}$ [5]. In addition, value 0 represents black pixels, while value 1 from (2) converted in number 255 represents white pixels. This circumstance is shown by the document text image fragment in Figure 2.

Figure 2. Document text image.

2.2. Test procedure

2.2.1. Multi-line text segmentation experiment

Algorithm quality examination consists of few text experiments representing the test procedure. In the first group of the experiments, text line segmentation quality is examined. These tests are significant because they are a prerequisite for obtaining the other text parameters. If segmentation experiment fails, then the examination of other features will be meaningless. Hence, its importance is critical. For this purpose, as the first experiment, a multi line text is used. Sample multi-line text with its skew angle parameter α is shown in Figure 3.

Figure 3. (a) Multi-line text referent line definition. (b) Multi-line text sample.**(a)**

*This is an example of the
multi-line text for testing
segmentation algorithm*

(b)

A number of existing text objects in a multi-line text image relate to the success of text segmentation. Hence, the less objects the better segmentation process, except the number may not be less than the number of text lines. As a quality measure, the root mean square error $RMSE_{seg}$ has been used. It is calculated as [7-9]:

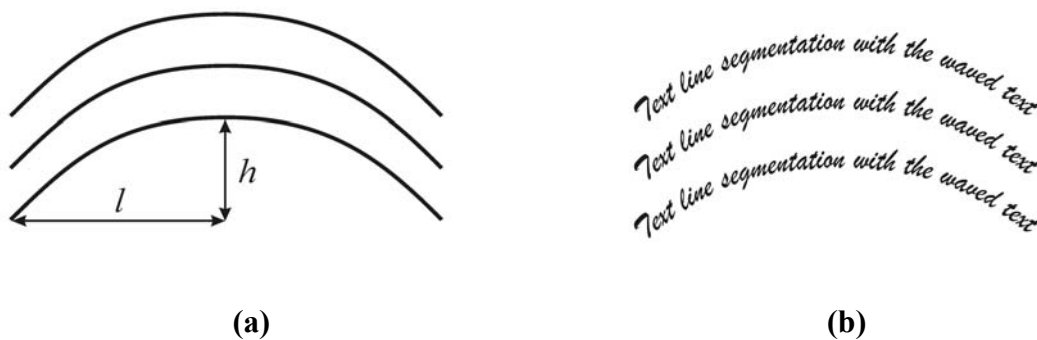
$$RMSE_{seg} = \sqrt{\frac{1}{P} \sum_{k=1}^P (O_{k,ref} - O_{k,est})^2} \quad (4)$$

where $k = 1, \dots, P$ is the number of examined text samples, $O_{k,ref}$ is the number of referent objects in text, *i.e.*, number of text lines, and $O_{k,est}$ is the number of objects obtained in the text by the applied algorithm.

2.2.2. Multi-line waved text segmentation experiment

The second text line segmentation experiment is a multi-line curved text one. Sample text is formed as a group of text lines using a curved reference line for its basis. The reference line is defined by the parameter $\varepsilon = h/l$. Typically, ε is used from the set $\{1/8, 1/6, 1/4, 1/3, \dots\}$. A sample multi-line curved text for the experiment is shown in Figure 4.

Figure 4. (a) Curved text referent line shape definition. (b) Curved multi-line text sample.



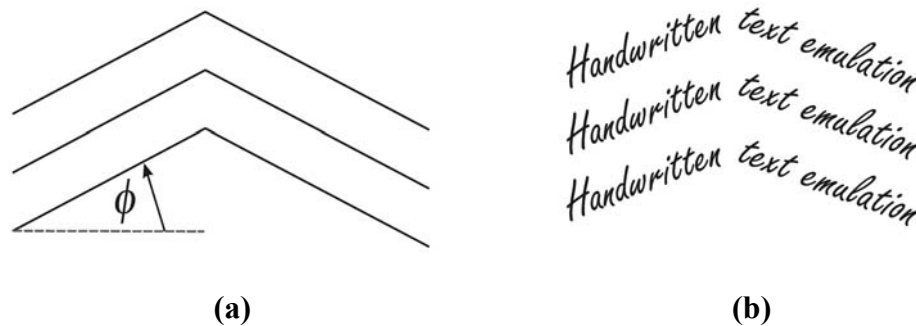
Like in the previous segmentation test, the number of existing text objects after the algorithm is applied relates to the text segmentation quality. Again, as a quality measure, the root mean square error $RMSE_{seg,wav}$ has been used. It is calculated as [7-9]:

$$RMSE_{seg,wav} = \sqrt{\frac{1}{R} \sum_{l=1}^R (O_{l,ref} - O_{l,est})^2} \quad (5)$$

where $l = 1, \dots, R$ is the number of examined text samples, $O_{l,ref}$ is the number of referent objects in text, *i.e.*, number of text lines, and $O_{l,est}$ is the number of objects obtained in the text by the applied algorithm.

2.2.3. Multi-line fractured text segmentation experiment

The last experiment in the first test group is a multi-line fractured text segmentation experiment. The sample text for this experiment is formed by using a reference fractured line as a basis. This fractured text reference line is defined by the slope angle ϕ , as a parameter. Typically, ϕ is used from the set $\{5^\circ, 10^\circ, 15^\circ, 20^\circ\}$. A sample multi-line fractured text for the last segmentation experiment is shown in Figure 5.

Figure 5. (a) Fractured text referent line slope definition. (b) Fractured multi-line text.

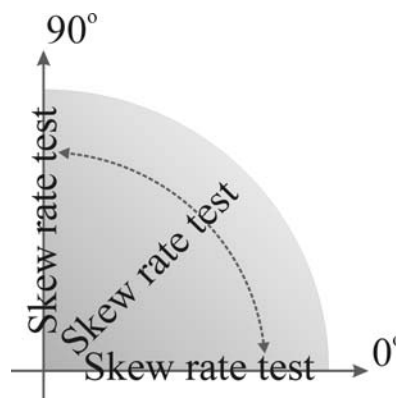
Again, the number of existing text objects relate to the text segmentation quality. The root mean square error $RMSE_{seg,frac}$ has been used. It is calculated as [7-9]:

$$RMSE_{seg,frac} = \sqrt{\frac{1}{Q} \sum_{m=1}^Q (O_{m,ref} - O_{m,est})^2} \quad (6)$$

where $m = 1, \dots, Q$ is the number of examined text samples, $O_{m,ref}$ is the number of referent objects in text, *i.e.*, number of text lines, and $O_{m,est}$ is the number of objects obtained in the text by the applied algorithm.

2.2.4. Skew rate text experiment

Further experiments belong in the second test group. The first of them, a skew rate test experiment, is mainly concerned with skew rate identification. It evaluates the algorithm's performance in the skew tracking domain. Although, this experiment is primarily based on printed text, it is good prerequisite for testing handwritten text as well. In this test, a sample printed text rotated from 0° to 90° in 5° steps around the x -axis is used. This is presented in Figure 6.

Figure 6. Sample text rotated up to 90° in 5° steps.

The reference line of the test sample text is represented by:

$$y = ax + b \quad (7)$$

After applying any algorithm to the sample text, reference text line estimation implies calculation of the average positions of only black pixels in every column of the document text image. It is calculated by [1,7,10]:

$$x_i = \frac{\sum_{j=1}^L y_j}{L} \quad i = 1, \dots, K \quad (8)$$

where x_i is the point position of calculated reference text line, i is the number of column position of the calculated reference text, y_j is the position of black pixel in column j and L is the sum of black pixel number in a specified column j of an image.

After calculation, an image matrix with only one black pixel per column is obtained. It defines the calculated i.e. estimated reference text line as well as text line skewness. This reference text line forms a continuous or discontinuous line partly or completely “representing” the reference text line. To achieve a continuous linear reference text line, the least squares method is used. The function is approximated by a first-degree polynomial is given by:

$$y = a'x + b' \quad (9)$$

Further, ndp represents the number of data points. It is used in the relation for calculating the slope a' , and the y -intercept b' as follows [8]:

$$a' = \frac{\sum y \sum xy - ndp \sum xy}{(\sum x)^2 - ndp \sum x^2} \quad (10)$$

and:

$$b' = \frac{\sum x \sum xy - \sum y \sum x^2}{(\sum x)^2 - ndp \sum x^2} \quad (11)$$

For algorithm approximation and evaluation, a quantity called relative error [9] is important. The reference line hit rate i.e., $RLHR$ incorporates this quantity. It is defined as [7,10]:

$$RLHR = 1 - \frac{\Delta\beta}{\beta} = - \frac{|\beta_{est} - \beta_{ref}|}{|\beta_{ref}|} \quad (12)$$

where β_{ref} is the arc tangent from the origin (7) i.e., a and β_{est} is the arc tangent from estimate (9), i.e., a' . Obviously, $RLHR$ is equal to 1– the relative error [9]. Now, the root mean square error $RMSE_{skew}$ is calculated by [7-10]:

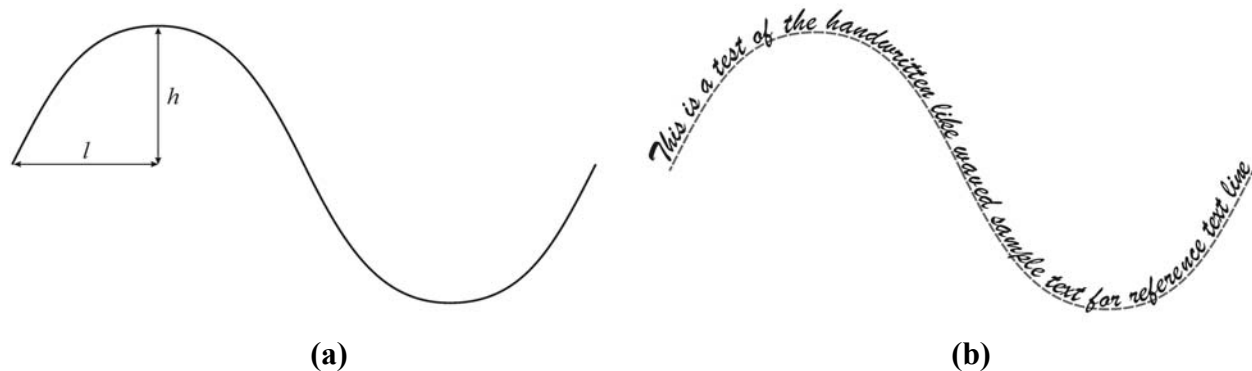
$$RMSE_{skew} = \sqrt{\frac{1}{S} \sum_{n=1}^S (O_{n,ref} - O_{n,est})^2} \quad (13)$$

where $n = 1, \dots, S$ is the number of examined text rotating angles up to 90° , $x_{n,ref}$ is $RLHR$ for β_{est} equal to β_{ref} , due to normalization equal to 1, and $x_{n,est}$ is $RLHR$.

2.2.5. Handwritten curved text experiment

The second experiment in this test group is primarily linked with handwritten text. Particularly, in this experiment a hypothetical reference text line is represented by a wavy line. This test examines the algorithm's capability to follow a wavy reference text line. This wavy text sample as well as its definition is given in Figure 7.

Figure 7. (a) Wavy text reference line slope and shape definition. (b) Wavy text sample.



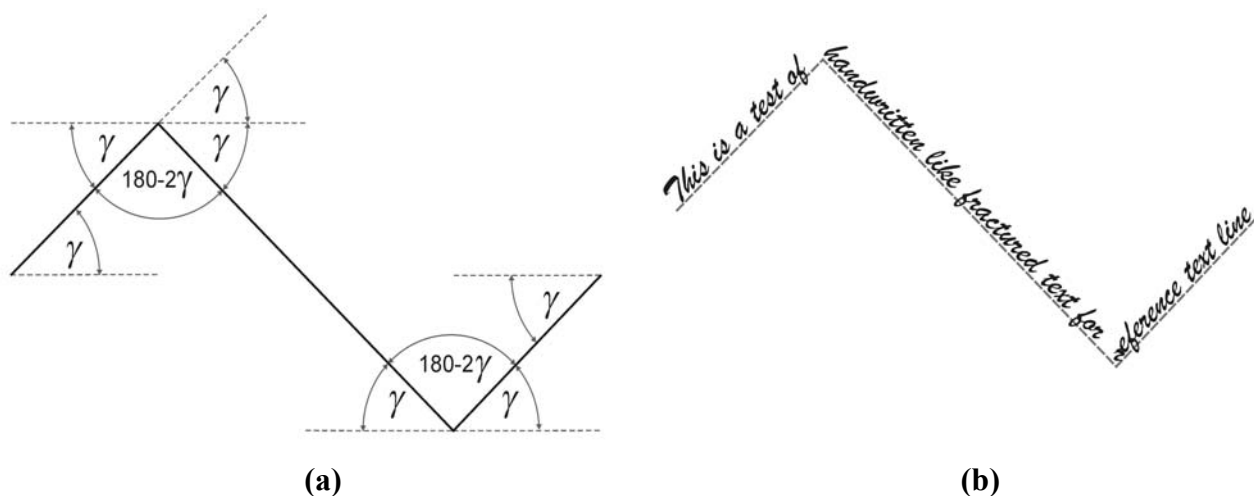
Different types of wavy text can be examined. As can be seen, it is completely defined by the ratio $\eta = h/w$ (see Figure 7a), where h represents height of the waved line, while w represents the half length of the wavy line. This experiment used the parameter set $\eta = \{1/8, 1/4, 1/2, 1\}$. Algorithm criteria quality and handwritten referent text line is measured and evaluated by the root mean square error $RMSE_{wav}$ calculated as [7-9]:

$$RMSE_{wav} = \sqrt{\frac{1}{T} \sum_{o=1}^T (O_{o,est} - O_{o,ref})^2} \quad (14)$$

where $o = 1, \dots, T$ is the number of examined text pixels, *i.e.*, columns of interest, $x_{o,ref}$ is pixel position of original referent text line in o -th column, and $x_{o,est}$ is pixel position of calculated, *i.e.*, estimated referent text line in o -th column.

2.2.6. Handwritten fractured text experiment

The next and the last experiment is also linked with handwritten text. A hypothetical reference text line is represented by a fractured line. This test examines the algorithm's ability to follow a fractured text line which represents abrupt changes in direction. The fractured text sample is given in Figure 8. In the figure, γ is the slope angle of the first, second and third part of the fractured text line. It can be observed that the second and third part of text line is rotated by an angle of 2γ from previous part of reference text line at once. Hence, this test example is a rather extreme one.

Figure 8. (a) Fractured text reference line slope definition. (b) Fractured text sample.

This experiment is typically performed for γ from 5° to 25° in 5° steps around the horizontal x -axis. Again, the evaluated reference text line is evaluated by root mean square error ($RMSE$) method. Further, $RMSE_{frac}$ is calculated as [7-9]:

$$RMSE_{frac} = \sqrt{\frac{1}{U} \sum_{p=1}^U (x_{p,est} - x_{p,ref})^2} \quad (15)$$

where $p = 1, \dots, U$ is number of examined text pixels, *i.e.*, columns of interest, $x_{p,ref}$ is pixel position of original referent text line in p -th column, and $x_{p,est}$ is pixel position of calculated, *i.e.*, estimated referent text line in p -th column.

2.2.7. Decision Making

Results obtained from different experiments during test procedure are inter-related. It should be noted that if the algorithm was examined and evaluated for text line segmentation as well as for text parameter extraction, then text line segmentation should be primary goal. Therefore, it is prerequisite for text parameters extraction such as reference text line and skew rate. This will be followed by the experiments for text parameter extraction and identification. Although, they are of second-rate importance compared to text line segmentation, their importance is evident at the next level of algorithm quality evaluation.

2.3. Combined test results

Combined test results represent merged results. Although, it is very similar to decision making, it is quite diverse from it. In fact, for the investigation of the algorithm under different parameters and restrictions, such merging of results leads to optimized parameter(s) value extraction. This way, an optimized subset of parameter(s) values is obtained. This process is invaluable for the algorithm evaluation as well as for obtaining any conclusions from it.

3. Test Example and Discussion

For the illustration, above test framework procedure will be examined using as an example the Gaussian isotropic kernel algorithm [7]. This algorithm will be just briefly explained. Its main task is expanding black pixel areas of text by scattering every black pixel in its neighborhood. This way, distinct areas that mutually separate text lines are established. Its primary purpose is joining only text elements from the same text line into the same distinct continuous areas. Gaussian probability function is taken as template that gives the probability of the random function. Consequently, it represents probability of the hypothetical expansion around every black pixel that represents a text element. Hence, around every black pixel, new pixels are non-uniformly dispersed. These new pixels have lower black intensity. Because the level of probability expansion relates to distance from black pixel, their intensity depends completely on their position *i.e.*, the distance from the original black pixel. Hence, these newly formed pixels are grayscale. Currently, document text image is represented by a grayscale image matrix. Thus, intensity pertains in level region $\{0, \dots, 255\}$. Hence, after applying Gaussian isotropic kernel, equal to $2K + 1$ in x -direction as well as in y -direction, text is scattered forming an enlarged area around it. Now, inside the kernel a “probability” sub area is formed using the radius 3σ , where σ represents standard deviation defining curve spread parameter. Converting all these pixels into black pixels as well as inverting image, forms the new black pixel expanded areas [7]. These areas are named boundary-growing areas.

The main purpose of the testing is optimization of the algorithm parameters. In our example, parameter of interest is K that defines kernel size. Further, the algorithm will be examined and evaluated. Firstly, the algorithm is examined by a multi-line text segmentation test. The multi-line text sample is skewed by an angle α (see Figure 3). Gaussian kernel size is defined by K value (in pixels), which is used as parameter. Because of the size of the letters, K is used from the set $\{5, 10, 15, 20, 25\}$. Obtained results are presented in Figures 9 and 10.

Figure 9. Number of the objects as a result of multi-line segmentation test.

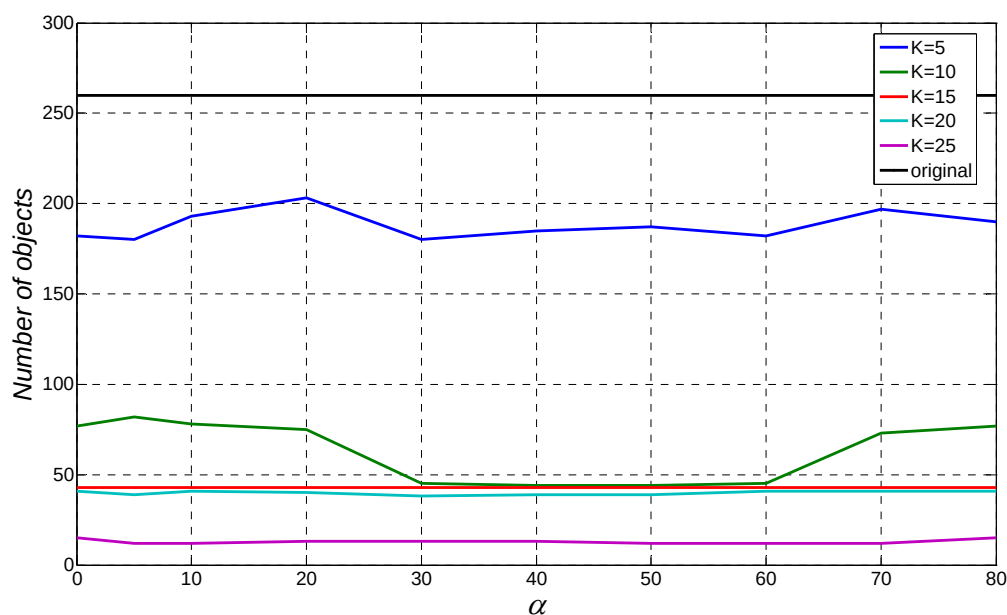
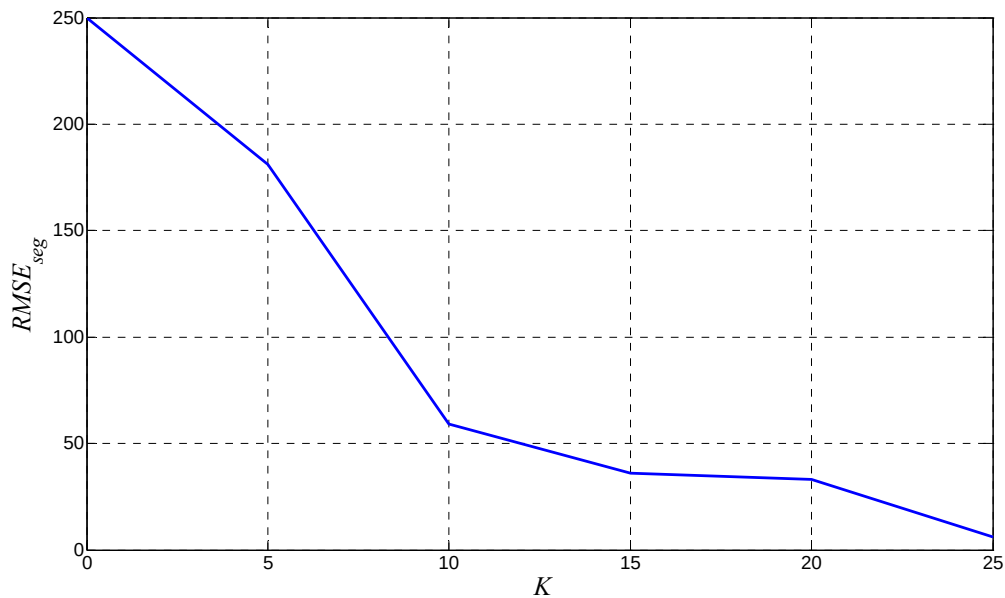
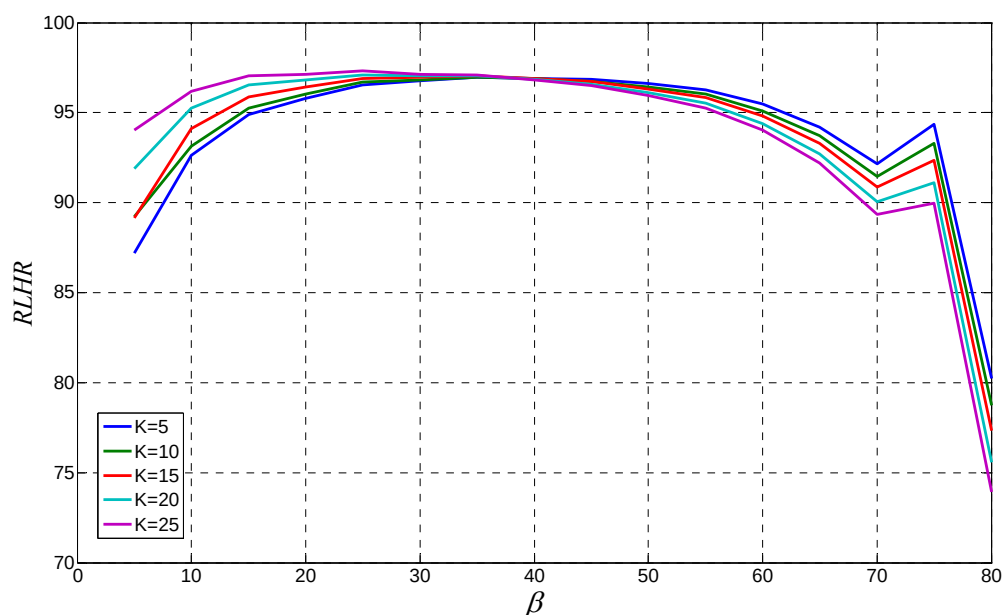


Figure 10. $RMSE_{seg}$ from the segmentation test results.

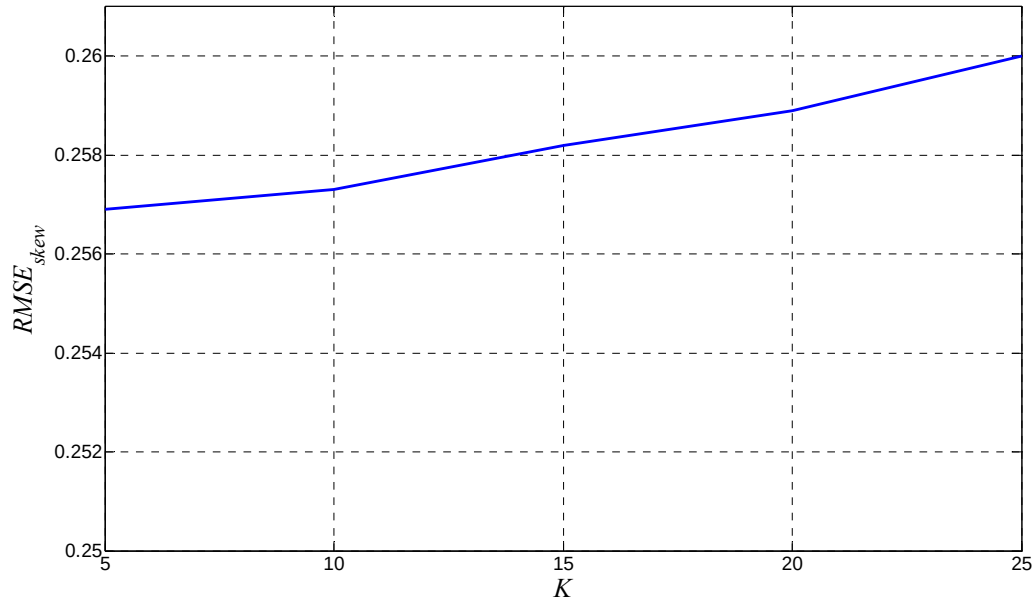
Although it is not part of the “relevant” measurement test, but just for the illustration purposes, the number of segmentation objects for different K is shown in Figure 9. Parameter $K = 0$ represents situation without an applied algorithm. Obviously, a bigger K value leads to better segmentation results due to stretching of the original text. Still, too big a K will join different text lines.

Similarly, the results presented in Figure 10 confirm that a big K , especially bigger than 15 leads to satisfactory segmentation results, while a small K , less than 10 are completely unacceptable. Hence, as a starting point, a K bigger than 15 is a good choice. Further, the algorithm is evaluated by a single line skew rate test. The obtained results are represented by the $RLHR$ value from (12). They are shown in Figure 11.

Figure 11. $RLHR$ for text rotated by angle β from 0° up to 80° .

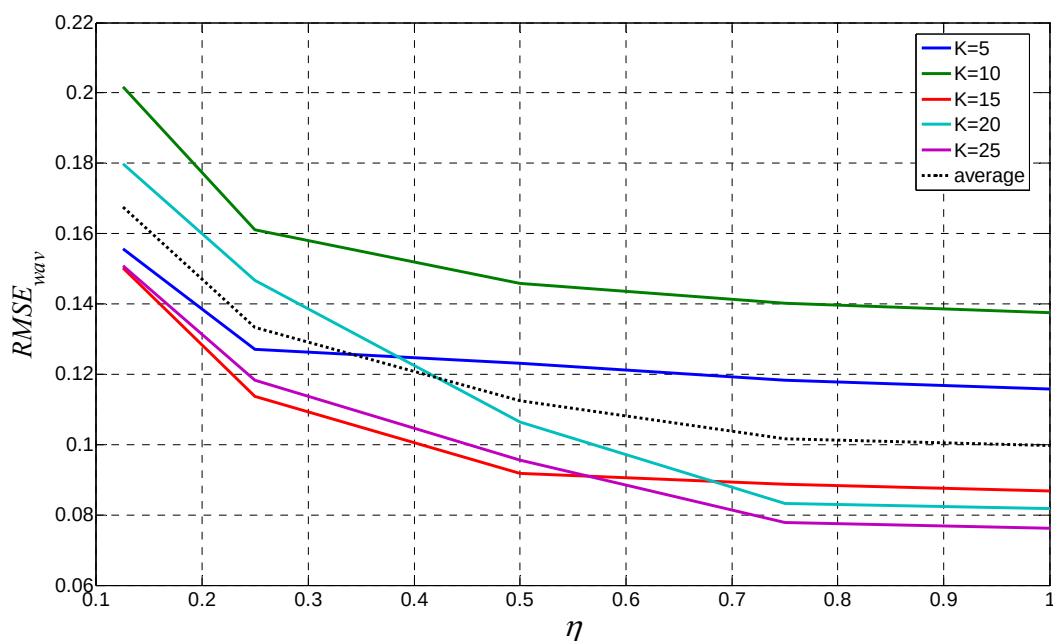
This time, a bigger K is not an advantage. Hence, a medium size K like $\{15, 20\}$ is the optimal value [7]. Scattering results from the previous test is represented by the $RMSE$ method in Figure 12.

Figure 12. $RMSE_{skew}$ from the single line skew test.



Bigger K leads to slightly bigger $RMSE_{skew}$. This is obvious because skewed reference line is more stretched by a bigger kernel which leads to slightly bigger scattering results. Finally, the algorithm is examined with a single line wavy text as well as a fractured text test. Algorithm criteria quality is measured by the $RMSE$ method. Results for the wavy text test are shown in Figure 13.

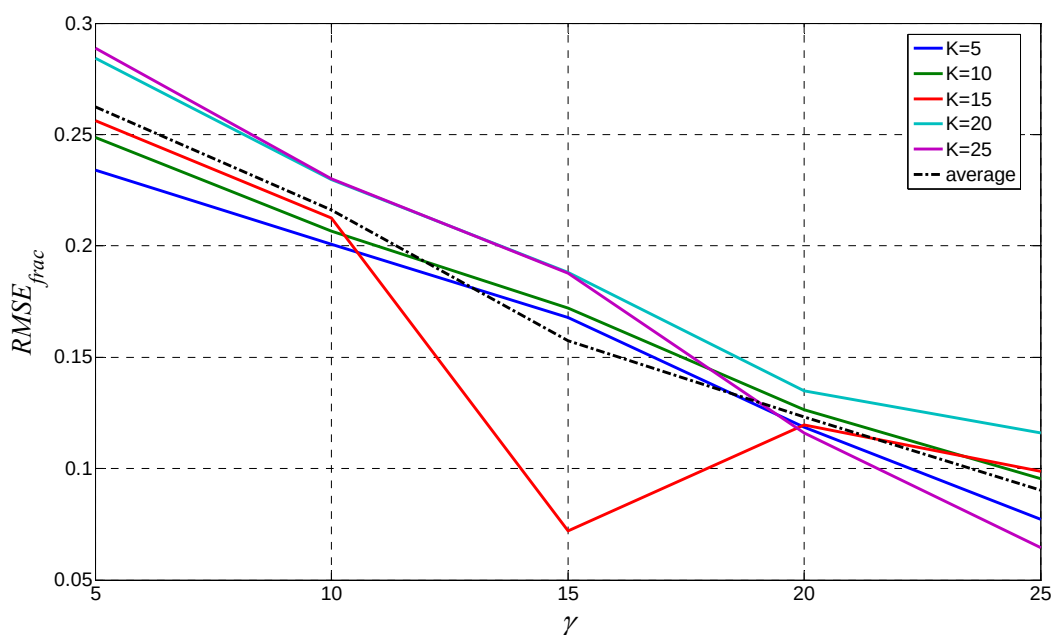
Figure 13. $RMSE_{wav}$ for the wavy text test.



From Figure 7a, the ratio $\eta = h/w$ is the severe element. Response to the sample wavy text is optimal for a K value chosen from the set $\{15, 20, 25\}$ that confirms their candidacy for the optimal values. Finally, results for the fractured text test are shown in Figure 14. The most promising $RMSE_{frac}$ value for this test is for $K = 15$. Hence, from all the above results, after intersection decision making, the optimal values for K are 15 or 20.

From the obtained and presented measurement results, it is obvious that the results of above tests are quite sufficient for the evaluation of the algorithm quality in the domain of the text segmentation and feature extraction. Hence, they can represent a basic test framework for the evaluation of the text line segmentation and text parameters extraction.

Figure 14. $RMSE_{frac}$ for the fractured text test.



4. Conclusions

The paper describes the proposal of a basic test framework for the evaluation of text feature extraction algorithms. All previous algorithm evaluation procedures were custom oriented. However, the proposed test framework is the first step toward testing generalization in the domain of document image processing algorithms. It consists of two groups of experiments. In the first and the most important group, text line segmentation experiments are included. These tests measure text line segmentation algorithm quality. They incorporate three various multi-line text experiments. Single line skew rate test belongs in the second group of the test framework. Its task is algorithm performance evaluation of the skew rate tracking success. Further tests in this group of experiments are primarily linked to “handwritten text”. They consist of single line wavy and fractured text tests. These tests examine an algorithm’s ability to follow wavy and fractured text reference lines. Results obtained from all test experiments are inter-related. Hence, after decision making and results merging, optimized values of the algorithm parameters are extracted. This way, an optimized subset of parameters values

is obtained. Hence, this process is invaluable for algorithm evaluation as well as for making any conclusions about it. In the end, its suitability for different types of letters and languages as well as its adaptability is a strong advantage.

References

1. Basu, S.; Chaudhuri, C.; Kundu, M.; Nasipuri, M.; Basu, D.K. Text Line Extraction from Multi-Skewed Handwritten Documents. *Pattern Recognition* **2006**, *40*, 1825–1839.
2. Amin, A.; Wu, S. Robust Skew Detection in mixed Text/Graphics Documents. In *Proceedings of International Conference on Document Analysis and Recognition (ICDAR)*, Seoul, Korea, August 29–September 1, 2005.
3. Likforman Sulem, L.; Zahour, A.; Taconet, B. Text Line Segmentation of Historical Documents: A Survey. *IJDAR* **2007**, *9*, 123–138.
4. Razak, Z.; Zulkiflee, K.; Idris, M.Y.I.; Tamil, E.M.; Noor, M.N.M.; Salleh, R.; Yaakob, M.; Yusof, Z.M.; Yaacob, M. Off-Line Handwriting Text Line Segmentation: A Review. *IJCSNS* **2008**, *8*, 12–20.
5. Gonzalez, R.C.; Woods, R.E. *Digital Image Processing*, 2nd ed.; Prentice-Hall: Bergen County, NJ, USA, 2002.
6. Otsu, N. A Threshold Selection Method from Gray-level Histograms. *IEEE Trans. Syst., Man, Cybern.* **1979**, *9*, 62–66.
7. Brodić, D.; Milivojević, Z. Using Anisotropic Gaussian Window for Printed and Handwritten Text Parameters Extraction. In *Proceedings of the 9th International Scientific Conference (UNITECH'09)*, Gabrovo, Bulgaria, September 10–14, 2009; Volume 1, pp. 453–460.
8. Bolstad, W.M. *Introduction to Bayesian Statistics*; John Wiley & Sons: Hoboken, NJ, USA, 2005; pp. 40–49, 115–116.
9. Terrell, G.R. *Mathematical Statistics: A Unified Introduction*; Springer-Verlag: New York, NY, USA, 1999; pp. 252–254.
10. Brodić, D.; Milivojević, Z. An Approach to Modification of Water Flow Algorithm for Segmentation and Text Parameters Extraction. In *Emerging Trends in Technological Innovation*; Camarinha-Matos, L.M., Pereira, P., Ribeiro, L., Eds.; Springer: Boston, MA, USA, 2010; IFIP AICT, Volume 314, pp. 324–331.

Appendix 1. Multi-line text segmentation test results.

α	Number of text objects					
	$K = 5$	$K = 10$	$K = 15$	$K = 20$	$K = 25$	$K = 0$
0°	182	77	43	41	15	260
5°	180	82	43	39	12	260
10°	193	78	43	41	12	260
20°	203	75	43	40	13	260
30°	180	45	43	38	13	260
40°	185	44	43	39	13	260

Appendix 1. Cont.

50°	187	44	43	39	12	260
60°	182	45	43	41	12	260
70°	197	73	43	41	12	260
80°	190	77	43	41	15	260
$RMSE_{seg}$	181.05	59.22	36.00	33.02	6.01	-

Appendix 2. Single line skew rate test results.

Reference line hit rate ($RLHR$)					
β	$K = 5$	$K = 10$	$K = 15$	$K = 20$	$K = 25$
5°	87.19%	89.24%	89.13%	91.88%	94.05%
10°	92.63%	93.14%	94.10%	95.24%	96.20%
15°	94.89%	95.26%	95.89%	96.53%	97.05%
20°	95.80%	96.02%	96.43%	96.81%	97.11%
25°	96.53%	96.68%	96.89%	97.10%	97.32%
30°	96.76%	96.81%	96.95%	97.06%	97.14%
35°	96.96%	96.97%	97.03%	97.04%	97.07%
40°	96.91%	96.90%	96.91%	96.84%	96.83%
45°	96.84%	96.75%	96.72%	96.59%	96.48%
50°	96.62%	96.44%	96.32%	96.11%	95.94%
55°	96.27%	96.02%	95.83%	95.53%	95.25%
60°	95.48%	95.09%	94.82%	94.40%	94.02%
65°	94.19%	93.72%	93.28%	92.70%	92.19%
70°	92.15%	91.47%	90.88%	90.06%	89.34%
75°	94.34%	93.30%	92.36%	91.12%	89.97%
80°	80.24%	78.73%	77.35%	75.57%	73.95%
$RMSE_{skew}$	0.2569	0.2573	0.2582	0.2589	0.2600

Appendix 3. Single line waved text test results.

$RMSE_{wav}$					
$\eta = h/w$	$K = 5$	$K = 10$	$K = 15$	$K = 20$	$K = 25$
1/8	0.1556	0.1271	0.123	0.1183	0.1158
1/4	0.2017	0.1611	0.1459	0.1402	0.1376
1/2	0.1502	0.1138	0.0919	0.0887	0.0869
3/4	0.1798	0.1467	0.1064	0.0832	0.0819
1	0.1509	0.1184	0.0955	0.0779	0.0761
$RMSE_{avg}$	0.1676	0.1334	0.1125	0.1017	0.0997

Appendix 4. Single line fractured text test results.

γ	$RMSE_{frac}$				
	$K = 5$	$K = 10$	$K = 15$	$K = 20$	$K = 25$
5°	0.234	0.2009	0.1677	0.1185	0.0772
10°	0.2487	0.2067	0.1721	0.1264	0.0956
15°	0.2563	0.2124	0.0719	0.1197	0.0986
20°	0.2842	0.2296	0.1882	0.1349	0.1160
25°	0.2888	0.2302	0.1876	0.116	0.0645
$RMSE_{avg}$	0.2624	0.2160	0.1575	0.1231	0.0904

© 2010 by the authors; licensee MDPI, Basel, Switzerland. This article is an Open Access article distributed under the terms and conditions of the Creative Commons Attribution license (<http://creativecommons.org/licenses/by/3.0/>).