

LAMP-TR-032
CAR-TR-906
CS-TR-3982

MDA 9049-6C-1250
January 1999

A Statistical, Nonparametric Methodology for Document Degradation Model Validation

Tapas Kanungo,¹ Robert Haralick,²
Henry Baird,³ Werner Stuezel,⁴ and David Madigan⁴

¹Center for Automation Research
University of Maryland
College Park, MD

³Xerox PARC
Palo Alto, CA

²Dept. of Electrical Engineering
University of Washington
Seattle, WA

⁴Dept. of Statistics
University of Washington
Seattle, WA

Abstract

Printing, photocopying and scanning processes degrade the image quality of a document. Statistical models of these degradation processes are crucial for document image understanding research. Models allow us to predict system performance; conduct controlled experiments to study the break-down points of the systems; create large multilingual data sets with groundtruth for training classifiers; design optimal noise removal algorithms; choose values for the free parameters of the algorithms; and so on. Although research in document understanding started many decades ago, only two document degradation models have been proposed thus far. Furthermore, no attempts have been made to statistically validate these models.

In this paper we present a statistical methodology that can be used to validate local degradation models. This method is based on a non-parametric, two-sample permutation test. Another standard statistical device — the power function — is then used to choose between algorithm variables such as distance functions. Since the validation and the power function procedures are independent of the model, they can be used to validate any other degradation model. A method for comparing any two models is also described. It uses p-values associated with the estimated models to select the model that is closer to the real world.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE JAN 1999		2. REPORT TYPE		3. DATES COVERED 00-01-1999 to 00-01-1999	
4. TITLE AND SUBTITLE A Statistical, Nonparametric Methodology for Document Degradation Model Validation				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Language and Media Processing Laboratory, Institute for Advanced Computer Studies, University of Maryland, College Park, MD, 20742-3275				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES The original document contains color images.					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES 30	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

LAMP-TR-032
CAR-TR-906
CS-TR-3982

MDA 9049-6C-1250
January 1999

**A Statistical, Nonparametric Methodology
for Document Degradation Model Validation**

Tapas Kanungo, Robert Haralick,
Henry Baird, Werner Stuehle, and David Madigan

A Statistical, Nonparametric Methodology for Document Degradation Model Validation

Tapas Kanungo,¹ Robert Haralick,²
Henry Baird,³ Werner Stuehle,⁴ and David Madigan⁴

¹Center for Automation Research
University of Maryland, College Park, MD

²Dept. of Electrical Engineering
University of Washington, Seattle, WA

³Xerox PARC, Palo Alto, CA

⁴Dept. of Statistics
University of Washington, Seattle, WA

Abstract

Printing, photocopying and scanning processes degrade the image quality of a document. Statistical models of these degradation processes are crucial for document image understanding research. Models allow us to predict system performance; conduct controlled experiments to study the break-down points of the systems; create large multilingual data sets with groundtruth for training classifiers; design optimal noise removal algorithms; choose values for the free parameters of the algorithms; and so on. Although research in document understanding started many decades ago, only two document degradation models have been proposed thus far. Furthermore, no attempts have been made to statistically validate these models.

In this paper we present a statistical methodology that can be used to validate local degradation models. This method is based on a non-parametric, two-sample permutation test. Another standard statistical device — the power function — is then used to choose between algorithm variables such as distance functions. Since the validation and the power function procedures are independent of the model, they can be used to validate any other degradation model. A method for comparing any two models is also described. It uses p-values associated with the estimated models to select the model that is closer to the real world.

This research was funded in part by the Department of Defense and the Army Research Laboratory under Contract MDA 9049-6C-1250.

1 Introduction

Printing, photocopying and scanning processes degrade the image quality of any document. Statistically valid models of these degradation processes can impact document image understanding research in many ways. Degradation models can be used to conduct controlled experiments to study the breakdown points of OCR systems; create large multilingual data sets with groundtruth for training classifiers; design optimal noise removal algorithms; choose values for the free parameters of the algorithms; predict OCR performance; and so on. Whereas research in document understanding started decades ago, only two document degradation models have been proposed thus far. Furthermore, no attempts have been made to statistically validate these models.

The current OCR evaluation methods rely on scanned documents, corresponding hand-entered ASCII groundtruth strings, and string matching algorithms that compare the groundtruth string against the OCR-generated string. The errors in the groundtruth are reduced by a process of cross-checking. This method is very expensive, laborious, and prone to errors. Furthermore, since the datasets are expensive, it is not possible to create large datasets that are representative of the variety of layout, font, and degradation levels seen in real-world documents. Despite these problems, various document databases with groundtruth have been created.

Our methodology for characterizing OCR algorithms is based on evaluating the algorithms on synthetically degraded documents. First, a wordprocessor is used to create an ideal document in any language, format or style. A bitmap version of this document is then created and degraded using a computer model of the real degradation process. This method has many advantages. First, since the ideal document is created using a wordprocessor, the groundtruth information associated with each character — location, identity, font type, etc. — is known without error. Second, the word processor can be used to reformat the documents (two columns, one column, various font types, sizes, etc.) to study the sensitivity of the OCR algorithm to these variables. Third, since the degradation model is under our control, we can create documents with varying levels of degradation and study how and where the OCR algorithm breaks down. Fourth, sample size is not a problem at all — any number of degraded samples can be created since all that needs to be done is to simulate another set of characters. In addition, there is no dearth of formatted documents — we create such documents daily, and so do academic journal publishers. Fifth, the model itself can be used in creating noise removal algorithms, training classifiers, choosing algorithm parameters, etc.

The main drawback of the above methodology is that it relies heavily on the simulation model being correct. That is, it assumes that the simulation model closely mimics reality. It is thus imperative that we validate the degradation model against real data. Only then can the simulations be used *in place of* real data. If the degradation model is not validated, results on the synthetically degraded documents should be used with caution, though they are still useful since they give *some* indication about the performance of the OCR algorithm.

In this paper we present a statistical methodology that can be used to validate the local degradation models. This method is based on a non-parametric, two-sample permutation test. Another standard statistical device — the power function — is then used

to choose between algorithm variables such as distance functions. Since the validation and the power function procedures are independent of the model, they can be used to validate any other degradation model. A method of comparing any two models is also described. It uses p-values associated with the estimated models to select the model that is closer to the real world.

In Section 2 we survey the related literature in the areas of degradation models, model validation, and statistical hypothesis testing and discuss their shortcomings. In Section 3 we describe our document degradation model for the local distortions that occur while printing, photocopying and scanning documents. The model is independent of the language in which the document is written. Our validation methodology is described in Section 4, and in Section 5 we apply it to datasets with known distributions to understand the performance of the permutation tests. In Section 6 we give experimental protocols and results of validation experiments on document images, and in Section 7 we present conclusions.

2 Related Literature

The earliest work on document degradation models is that of Baird [2, 3, 4]. Unfortunately, his degradation model is not validated. Furthermore, his paper advocates the use of isolated, synthetically degraded characters. Thus the degradation due to merging of neighboring characters is not reflected in his model. In addition, the uni-gram and bi-gram occurrence probabilities of characters in real-world text are not reflected in isolated-character experiments.

In contrast, our document degradation model, which is described in Section 3, advocates the use of complete documents for generating synthetically degraded characters. It thus takes into account the degradations arising due to merging of characters, the occurrence probabilities of individual characters, and the variability in the layout structure of the documents. The pixel degradations themselves are based on a local morphological model, which models the final spatial characteristics of the degradation process rather than the underlying physical process.

To the best of our knowledge, the only other work on validation of document degradation models is that of Nagy and Lopresti [23, 21, 22]. They are of the opinion that a degradation model is valid if the OCR confusion matrices that result from synthetically degraded documents are similar to the OCR confusion matrices produced from real documents. Unfortunately, this methodology validates the model-OCR combination and not the model itself. For instance, if the OCR system automatically filters noise, their validation process will not detect any difference between the real documents and the synthetically degraded documents even if the degradation process adds noise to the document. Furthermore, although they treat the OCR as a black box, the OCR algorithm itself has many parameters that can greatly influence the decisions of the validation procedure. Another drawback of their approach is that they do not indicate how their validation procedure can be compared to other validation procedures.

Our validation method, on the other hand, reduces the problem of model validation to a nonparametric statistical hypothesis testing problem, which is a well studied and accepted method in statistics [8, 7]. In addition, we use simple character distance functions

for the validation procedure, instead of entire OCR systems. Although the validation process now depends on these distance functions, they are much simpler than OCR black boxes. Finally, we provide a technique for comparing our validation method with other validation methods. This comparison procedure is based on “power functions,” which again are standard statistical devices for comparing hypothesis testing procedures.

3 A Document Degradation Model

In this section we describe a document degradation model for local distortions that are introduced during the printing, photocopying and scanning processes. A model for the perspective and illumination distortions that get introduced when we photocopy or scan thick bound books is described in [15, 16, 12].

Our local document degradation model accounts for (i) pixel inversion (from foreground to background and vice versa) that occurs independently at each pixel due to light intensity fluctuations, sensitivity of the sensors, and the thresholding level, and (ii) blurring that occurs due to the point-spread function of the scanner optical system. We model the pixel-flipping probability of a background pixel as an exponential function of its distance from the nearest boundary pixel. The parameter α_0 is the initial value for the exponential and the decay speed of the exponential is controlled by the parameter α . The foreground and background 4-neighbor distance are computed using a standard distance transform algorithm (see [5]). The flipping probabilities of the foreground pixels are similarly controlled by β_0 and β . The parameter η is the constant probability of flipping for all pixels. Finally, the last parameter k , which is the size of the disk used in the morphological closing operation, accounts for the correlation introduced by the point-spread function of the optical system.

The degradation model thus has six parameters: $\Theta = (\eta, \alpha_0, \alpha, \beta_0, \beta, k)^t$. These parameters are used to degrade an ideal binary image as follows:

1. Compute the distance d of each pixel from the character boundary.
2. Flip each foreground pixel with probability $p(0|1, d, \alpha_0, \alpha) = \alpha_0 e^{-\alpha d^2} + \eta$.
3. Flip each background pixel with probability $p(1|0, d, \beta_0, \beta) = \beta_0 e^{-\beta d^2} + \eta$.
4. Perform a morphological closing operation with a disk structuring element of diameter k .

Software for simulating noisy documents using the above degradation model is available from the University of Washington English Document Database I, and the model has also been described in the literature [16, 15]. The application of the various steps of the model is illustrated in Figure 1. In Figure 1(a) we show an ideal character. The distance transform of the foreground of (a) is shown in (b). The brighter pixels are further away from the pixel boundary. The distance transform of the background is shown in Figure 1(d). The ideal image after its pixels have been flipped according to the model is shown in (c). The final image after the closing operation is shown in (e).

The procedure described above works on bit-mapped images. Since there is no restriction on the size of the image that can be degraded, or the language of the written text, an

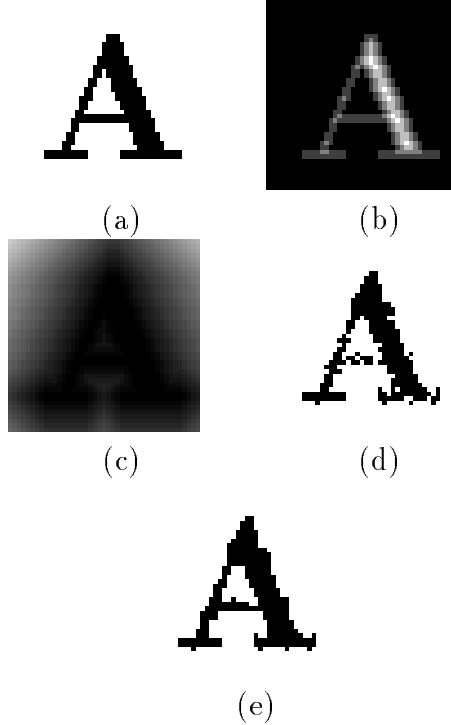


Figure 1: Local document degradation model: (a) Ideal noise-free character; (b) Distance transform of the foreground; (c) Distance transform of the background; (d) Result of the random pixel-flipping process (the probability of a pixel flipping is $p(0|d, \beta, f) = p(1|d, \alpha, b) = \alpha_0 e^{-\alpha d^2}$; here $\alpha = \beta = 2$, $\alpha_0 = \beta_0 = 1$); (e) Morphological closing of the result in (d) by a 2×2 binary structuring element.

entire document page image can be degraded using this model. In fact, since typesetting is under the experimenter's control, the same text can be re-formatted in various styles (single column, multiple column, report, book, etc.), font types (Roman, Helvetica, etc), and font sizes (9pt, 10pt, 12pt, etc.). Thus the performance of any character recognition system can be studied by providing as input the same (or different) text formatted in various styles with varied but controlled degradation.

We now show examples where we degrade complete document pages using our degradation model. In Figure 2(a), we show an ideal document formatted in $\text{\LaTeX}$ using the IEEE Transactions typesetting style. In Figure 2(b) we show a degraded version of the document in Figure 2(a).

The noise-free documents are typeset using the $\text{\LaTeX}$ formatting system [20, 19]. The ASCII files containing the text and the $\text{\LaTeX}$ typesetting information are then converted into a device-independent format (DVI) using $\text{\LaTeX}$. A software program called DVI2TIFF — which is a modified version of a DVI file previewer called XDVI [24] — is run to produce one bit/pixel binary images in TIFF format from the DVI files. Besides producing the binary images of the documents, DVI2TIFF also produces the groundtruth information regarding each character in the document image.

The local document degradation model is another software program called DDM. This program takes as input an ideal binary document image in TIFF format, and a file

This Is A Sample File Using The ‘IEEEtran.sty’, To Help You Estimate Your Page Count And Facilitate Input-Processing Of Your Compuscript

ERB, Woody, Pheff, Bont, Tranman, IP, Dalton, Christine and OOO

Abstract—The theoretical analysis and derivation of artificial neural systems consist essentially of manipulating symbolic mathematical objects according to certain mathematical and biological knowledge. A simple observation has been made that this work can be done more efficiently with computer assistance by using and extending methods and systems of symbolic computation. In this paper, after presenting the mathematical characteristics of neural systems and a brief review on Lyapunov stability theory, we present some features and capabilities of existing systems and our extension for manipulating objects occurring in the analysis of neural systems. Then, some strategies and a toolkit developed in MACSYMA for computer aided analysis and derivation are described. A concrete example is given to demonstrate the derivation of a hybrid neural system, i.e. a system which in its learning rule combines elements of supervised and unsupervised learning. The future work and directions on this topic are indicated.

Keywords—CA system, computer aided analysis and derivation, Lyapunov function, neural system, symbolic computation.

I. INTRODUCTION

SINCE the early 1940s a large number of artificial neural systems have been proposed by neural scientists. The dynamical behavior of these systems may be mathematically described by sets of coupled equations like differential equations for formal neurons with graded response. The investigation of essential features of neural systems such as stability and adaptation depends strongly upon the state of the mathematical theory to be applied and on a concrete and efficient analysis of dynamical equations. Unlike abstract theoretical research in which the mathematical objects adopted are frequently assumed to be of certain canonical form, the neurodynamics is usually complicated due to various biological facts which should be taken account of to a degree as large as possible. Consequently, this makes the analysis and derivation very complex, sometimes to an extent which is beyond human capacity, and the traditional methods and tools of mathematics are not always sufficient. It is therefore proposed in [19] to use and extend the methods and software systems of symbolic computation for handling, analyzing and constructing neurodynamics and its related objects. The present paper is the continuation of our work in this direction. The attempt is to demonstrate how symbolic computation can be applied to aid the analysis and derivation of neural systems.

This would be where the author affiliation would be, together with the IEEE log Number, like so:
The research of G. Wang is supported by a grant from SIEMENS AG Munich. He is with the Research Institute for Symbolic Computation, Johannes Kepler University, 4040 Linz, Austria.
B. Schirmann is with the Corporate Research and Development, ZPE IS INF 2, Siemens AG, 8000 München 83, Germany.

In contrast to the approximative character of numerical calculations, symbolic computation treats objects with semantics like functions, formulae and programs. A variety of software systems for performing symbolic computation have been developed for research and applications in natural and technical sciences. However, the existing systems cannot be directly used for the analysis and derivation of neural systems as the operations on the occurring objects, particularly those involving an unspecified number of arguments like indefinite summations, have not yet been taken into account. To achieve our goal, some rules for differentiating and integrating indefinite summations with respect to indexed variables were proposed [20]. A toolkit has been designed and implemented in MACSYMA for manipulating these objects occurring in the analysis and derivation of neural systems [21].

In the next section, we introduce the general method and techniques for the stability analysis of artificial neural systems. The role of symbolic computation for representing and manipulating the objects concerning neural systems is discussed in Section III. In Section IV we present some strategies for using computer algebra (CA) systems and their extension to analyse the stability of neural systems and to derive novel stable systems. A brief description of a toolkit developed in MACSYMA is also provided. A concrete example is given in Section V to illustrate the derivation of a hybrid model by our toolkit. Section VI contains a discussion on future developments. The paper is closed with a brief summary.

II. STABILITY ANALYSIS OF NEURAL SYSTEMS

Consider artificial neural systems which are described by coupled systems of differential equations of the form

$$\dot{z} = F(z, w, K) \quad (1)$$

and

$$\dot{w} = G(z, w, K) \quad (2)$$

where $z = (z_1(t), \dots, z_n(t))$ is the activation state vector, $w = (w_1(t))$ is the weight matrix of dimension $n \times m$, n is the number of nodes and K is an external time-independent pattern vector. Such systems of differential equations which describe the neural model will occasionally be named *neurodynamics*.

Once a neural model is proposed, its main features are represented by its dynamic behavior. The adaptability of

This Is A Sample File Using The ‘IEEEtran.sty’, To Help You Estimate Your Page Count And Facilitate Input-Processing Of Your Compuscript

ERB, Woody, Pheff, Bont, Tranman, IP, Dalton, Christine and OOO

Abstract—The theoretical analysis and derivation of artificial neural systems consist essentially of manipulating symbolic mathematical objects according to certain mathematical and biological knowledge. A simple observation has been made that this work can be done more efficiently with computer assistance by using and extending methods and systems of symbolic computation. In this paper, after presenting the mathematical characteristics of neural systems and a brief review on Lyapunov stability theory, we present some features and capabilities of existing systems and our extension for manipulating objects occurring in the analysis of neural systems. Then, some strategies and a toolkit developed in MACSYMA for computer aided analysis and derivation are described. A concrete example is given to demonstrate the derivation of a hybrid neural system, i.e. a system which in its learning rule combines elements of supervised and unsupervised learning. The future work and directions on this topic are indicated.

Keywords—CA system, computer aided analysis and derivation, Lyapunov function, neural system, symbolic computation.

I. INTRODUCTION

SINCE the early 1940s a large number of artificial neural systems have been proposed by neural scientists. The dynamical behavior of these systems may be mathematically described by sets of coupled equations like differential equations for formal neurons with graded response. The investigation of essential features of neural systems such as stability and adaptation depends strongly upon the state of the mathematical theory to be applied and on a concrete and efficient analysis of dynamical equations. Unlike abstract theoretical research in which the mathematical objects adopted are frequently assumed to be of certain canonical form, the neurodynamics is usually complicated due to various biological facts which should be taken account of to a degree as large as possible. Consequently, this makes the analysis and derivation very complex, sometimes to an extent which is beyond human capacity, and the traditional methods and tools of mathematics are not always sufficient. It is therefore proposed in [19] to use and extend the methods and software systems of symbolic computation for handling, analyzing and constructing neurodynamics and its related objects. The present paper is the continuation of our work in this direction. The attempt is to demonstrate how symbolic computation can be applied to aid the analysis and derivation of neural systems.

This would be where the author affiliation would be, together with the IEEE log Number, like so:
The research of G. Wang is supported by a grant from SIEMENS AG Munich. He is with the Research Institute for Symbolic Computation, Johannes Kepler University, 4040 Linz, Austria.
B. Schirmann is with the Corporate Research and Development, ZPE IS INF 2, Siemens AG, 8000 München 83, Germany.

In contrast to the approximative character of numerical calculations, symbolic computation treats objects with semantics like functions, formulae and programs. A variety of software systems for performing symbolic computation have been developed for research and applications in natural and technical sciences. However, the existing systems cannot be directly used for the analysis and derivation of neural systems as the operations on the occurring objects, particularly those involving an unspecified number of arguments like indefinite summations, have not yet been taken into account. To achieve our goal, some rules for differentiating and integrating indefinite summations with respect to indexed variables were proposed [20]. A toolkit has been designed and implemented in MACSYMA for manipulating these objects occurring in the analysis and derivation of neural systems [21].

In the next section, we introduce the general method and techniques for the stability analysis of artificial neural systems. The role of symbolic computation for representing and manipulating the objects concerning neural systems is discussed in Section III. In Section IV we present some strategies for using computer algebra (CA) systems and their extension to analyse the stability of neural systems and to derive novel stable systems. A brief description of a toolkit developed in MACSYMA is also provided. A concrete example is given in Section V to illustrate the derivation of a hybrid model by our toolkit. Section VI contains a discussion on future developments. The paper is closed with a brief summary.

II. STABILITY ANALYSIS OF NEURAL SYSTEMS

Consider artificial neural systems which are described by coupled systems of differential equations of the form

$$\dot{z} = F(z, w, K) \quad (1)$$

and

$$\dot{w} = G(z, w, K) \quad (2)$$

where $z = (z_1(t), \dots, z_n(t))$ is the activation state vector, $w = (w_1(t))$ is the weight matrix of dimension $n \times m$, n is the number of nodes and K is an external time-independent pattern vector. Such systems of differential equations which describe the neural model will occasionally be named *neurodynamics*.

Once a neural model is proposed, its main features are represented by its dynamic behavior. The adaptability of

1

(a)

1

(b)

Figure 2: (a) An ideal document page typeset using L^AT_EX and IEEE Transactions typesetting style. (b) A synthetically degraded version of the document in (a).

containing the degradation model parameter values, and produces the binary degraded images in TIFF format.

Both programs — DVI2TIFF and DDM — are implemented using the C language and have been tested on SUN and IBM machines running the UNIX operating system. The software is available on the UW CD-ROM-1 [9].

4 Model Validation and Parameter Estimation

4.1 Statistical Problem Definition

In this section we formulate the degradation model validation problem as a statistical problem. Although degradation of the document is over the entire page, the degradation process itself is local. That is, degradation in one region does not influence the degradation process in another sufficiently distant region. More precisely, the degradation at a pixel is influenced only by pixels within a local neighborhood. Thus, one way

to characterize the degradation process is to study the degradation of local patterns. Since the most common patterns that occur on a document page are characters, we statistically characterize the degradation of individual characters on the page and use this characterization to estimate the parameters of a degradation model that produces similar degradations.

Assume that a scanned character is represented by a 30×30 matrix of zeros and ones. This matrix can be represented as a 1000×1 vector x ($30 \times 30 \approx 1000$). Let B be the space of $D = 1000$ -dimensional binary vectors, that is, $B = \{0, 1\}^D$. Now, let $x_1, x_2, \dots, x_N \in B$ be independent and identically distributed D -dimensional vectors representing instances of degraded characters produced from the same class ω . That is, each x_i is a degraded character that is produced from the same ideal pattern ω (say the ideal character ‘e’) and the same degradation process. The validation problem we need to address is:

Suppose we are given a set of *real* degraded instances $x_1, \dots, x_N \in B$ of the pattern ω and another set of *synthetic* degraded instances $y_1, \dots, y_M \in B$ of the pattern ω . Test the null hypothesis that $y_1, \dots, y_M$ and $x_1, \dots, x_N$, are samples taken from the same underlying population, to a specified significance level ϵ .

In our case D is large, typically on the order of 1000. Thus the number of possible x_i ’s is 2^{1000} , which is approximately equal to 10^{300} — a dauntingly large number. The available sample size, N , is typically on the order of 1000. Thus the x_i ’s occupy the space B extremely sparsely, which implies that the density function cannot be estimated reliably from the sample. This fact prohibits us from performing any standard statistical test based on density estimates. In the next section we describe a non-parametric method that overcomes this problem.

4.2 Permutation Tests and Model Validation

In this section we describe a nonparametric validation procedure that can be used to statistically validate any document degradation model. Suppose we are given a set of real degraded characters $X = \{x_1, x_2, \dots, x_N\}$, and another set of artificially degraded characters $Y = \{y_1, y_2, \dots, y_M\}$ that were created by perturbing an ideal character with a document degradation model. We can assume that the characters x_i and y_i are binary matrices of size (approximately) 30×30 . Note that every x_i and y_i can be of different size because the scanned characters can be of different sizes. The question that needs to be addressed is whether or not the x_i ’s and the y_i ’s come from the same underlying population. At this point it does not matter where the x_i ’s and the y_i ’s came from; they could both be synthetically generated, or both be real instances, or one of them could be synthetic and the other real. A statistical hypothesis test can be performed to test the null hypothesis that the underlying population distributions of the x_i ’s and y_i ’s are the same.

Standard parametric hypothesis testing procedures are not usable for our problem because the sizes of the binary matrices x_i and y_i are not fixed. Furthermore, the size of the space to which they belong is very large (approximately 2^{900} if we assume each

character to be of size 30×30) and so while in principle it is possible to estimate the density function, in practice it is not possible to do so because of the small sample size. Instead, we now describe a *nonparametric permutation test* (see [8, 7]) that performs this hypothesis test.

1. Given (i) the real data $X = \{x_1, x_2, \dots, x_N\}$, (ii) the synthetic data $Y = \{y_1, y_2, \dots, y_M\}$, (iii) a distance function $\rho(X, Y)$ on sets, (iv) a distance function $\delta(x, y)$ on characters, and (v) the size ϵ of the test (i.e. misdetection rate $= \epsilon$).
2. Compute $d_0 = \rho(X, Y)$.
3. Create a new sample $Z = \{x_1, \dots, x_N, y_1, \dots, y_M\}$. Thus Z has $N + M$ elements labeled $z_i, i = 1, \dots, N + M$.
4. Randomly permute (reorder) Z .
5. Partition the set Z into two sets X' and Y' where $X' = \{z_1, \dots, z_N\}$ and $Y' = \{z_{N+1}, \dots, z_{N+M}\}$.
6. Compute $d_i = \rho(X', Y')$.
7. Repeat steps 4, 5 and 6 K times to get K distances $d_1, \dots, d_K$.
8. Compute the empirical distribution of the d_i 's: $P(d \geq v) = \#\{k | d_k \geq v\} / K$
9. Compute the p-value: $\epsilon_0 = P(d \geq d_0)$.
10. Reject the null hypothesis that the two samples come from the same population if $\epsilon_0 < \epsilon$.

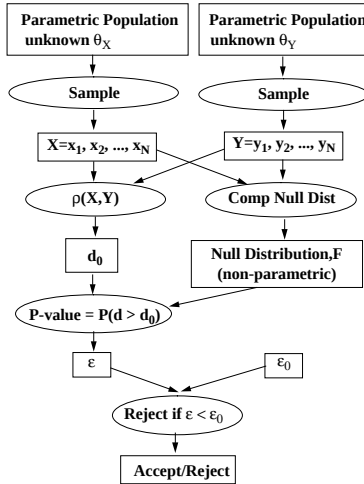


Figure 3: Here we show how the nonparametric test works when the two samples X and Y are from arbitrary distributions. For our problem, x_i and y_i are binary characters. In this case the null distribution cannot be determined theoretically.

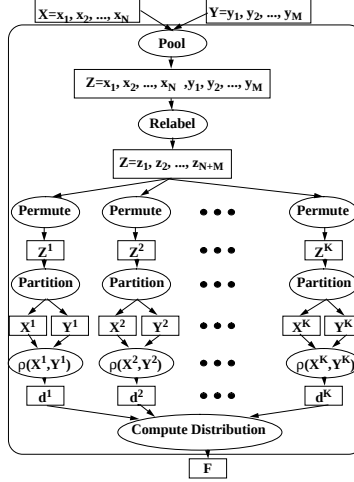


Figure 4: This figure shows the permutation procedure for computing the null distributions.

The above procedure computes the null distribution of the distance function $\rho(X, Y)$ nonparametrically. In a standard parametric hypothesis-testing procedure, the forms of the distributions of x and y are known (usually assumed to be Gaussian) and so the null distribution of the test statistic $\rho(X, Y)$ is known. In contrast, the permutation test does not make any prior assumption regarding the distributions of x and y . Instead, an empirical null distribution is created by randomly permuting the data set Z and creating a histogram of computed test statistics (d_i 's).

By design, the size of the test, ϵ , is fixed. Thus, irrespective of the distance function $\rho(X, Y)$, the percentage of time that the validation procedure rejects a true null hypothesis that the two samples are from the same underlying population is ϵ . In other words, the probability of misdetection is ϵ . What is not fixed is the probability of false alarm, γ , which is the probability that the procedure claims that X and Y come from the same underlying population when, in fact, they come from different underlying populations. Although the use of various distance functions for ρ and δ gives rise to the same probability of misdetection ϵ , each has a different probability of false alarm γ ,

It is important to note that if two samples X and Y pass the validation procedure, this does not mean that we accept the null hypothesis. Rather, it means that we do not have enough statistical evidence to reject the null hypothesis. Nevertheless, when we reject a null hypothesis, this *does* mean that we have enough statistical evidence to reject it.

4.3 Power Functions

Let us assume that the x_i 's are distributed as $F(\theta_X)$ and the y_i 's are distributed as $F(\theta_Y)$, where θ_X and θ_Y are the parameters of the distributions. Let the null hypothesis H_N and the alternate hypothesis H_A be

$$H_N : \quad \theta_X = \theta_Y \quad (1)$$

$$H_A : \quad \theta_X \neq \theta_Y \quad (2)$$

The size of the test, ϵ , is fixed by the algorithm designer and is given as

$$\epsilon = P(H_A | H_N \text{ is true}) . \quad (3)$$

The plot of 1 minus the probability of false alarm as a function of θ is the *power function*. Thus, if we fix the distribution parameter of the x_i 's at $\theta_X = \theta_0$, and vary the distribution parameter value $\theta_Y = \theta$ for the y_i 's, the power function is denoted by $\gamma_{\theta_0}(\theta)$, and is given by

$$\gamma_{\theta_0}(\theta) = P(H_A | \theta_X = \theta_0 \text{ and } \theta_Y = \theta) . \quad (4)$$

Thus $1 - \gamma_{\theta_0}(\theta)$ is the probability of false alarm. The power function should have a minimum at $\theta_X = \theta_Y = \theta_0$, with $\gamma_{\theta_0}(\theta_0) = \epsilon$, and should increase on either side and go up to 1 when $\theta_Y = \theta$ is very far from θ_0 .

Let us say there are two validation schemes A and B with test size ϵ and power functions $\gamma_{\theta_0}^A(\theta)$ and $\gamma_{\theta_0}^B(\theta)$. Since the misdetection probability ϵ is the same for both schemes, A is better than B if the false alarm rate of A is less than the false alarm rate of B . That is, A is better than B if $1 - \gamma_{\theta_0}^A(\theta) < 1 - \gamma_{\theta_0}^B(\theta)$ or $\gamma_{\theta_0}^A(\theta) > \gamma_{\theta_0}^B(\theta)$. If this relation is true for all values of θ , the procedure A is said to be uniformly more powerful than B . That is, scheme A is better than scheme B if the power function plot of A is above the power function plot of B for all values of θ . Generalizing, if there are many validation schemes, the one whose power function is above all other power functions is the best scheme. If the power functions intersect, there is no clear winner; for some regions in the parameter space one scheme is better while in other regions the other scheme is better.

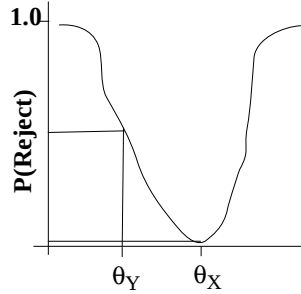


Figure 5: The true parameter of the sample X is Θ_X . The parameter Θ_Y of the sample Y is updated and the corresponding probability of the test rejecting the null hypothesis that X and Y are from the same underlying distribution is plotted. The resulting curve is the power function.

For a given validation scheme, if we increase the sample sizes N and M , the power function changes and the new power function is higher than the old power function, and so by definition is more powerful. Thus the sensitivity, i.e., the width of the notch at the minimum, is a function of the sample sizes N and M . When the sample size is small, the notch is broader, and when the sample size is large, the notch is sharper. This fact is used in deciding what sample size should be used for the test: choose the sample size such that the desired probability of false alarm is attained when the parameters θ_X and θ_Y differ by a specified amount $\Delta\theta$.

Finally, our validation scheme described in the previous section is dependent on two distance functions ρ and δ . Thus, each choice of ρ and δ gives rise to a different power function. The combination that produces the highest power function is the best choice. See [1] for details on power functions.

4.4 Distance Functions, Outliers, and Robust Statistics

Various distance functions $\rho(X, Y)$ can be used for computing the distance between the sets of characters X and Y . We use the following symmetric distance functions for ρ .

Mean Nearest Neighbor Distance:

$$\rho(X, Y) = \rho_{Mean}(X, Y) = \frac{\rho_{Mean}(Y; X) + \rho_{Mean}(X; Y)}{N + M}$$

where

$$\begin{aligned}\rho_{Mean}(Y; X) &= \sum_{x \in X} \left(\min_{y \in Y} \delta(x, y) \right) \\ \rho_{Mean}(X; Y) &= \sum_{y \in Y} \left(\min_{x \in X} \delta(x, y) \right).\end{aligned}$$

Trimmed Mean Nearest Neighbor Distance:

$$\rho(X, Y) = \rho_{Trim}(X, Y) = \frac{\rho_{Trim}(Y; X) + \rho_{Trim}(X; Y)}{2}$$

where

$$\begin{aligned}\rho_{Trim}(Y; X) &= \text{Trim}_{x \in X} \left(\min_{y \in Y} \delta(x, y) \right) \\ \rho_{Trim}(X; Y) &= \text{Trim}_{y \in Y} \left(\min_{x \in X} \delta(x, y) \right).\end{aligned}$$

Here the *Trim* function accepts as input a set of real numbers, orders them, and then discards the top and bottom 10% and returns the mean of the remaining 80%.

Median Nearest Neighbor Distance:

$$\rho(X, Y) = \rho_{Med}(X, Y) = (\rho_{Med}(Y; X) + \rho_{Med}(X; Y))/2$$

where

$$\begin{aligned}\rho_{Med}(Y; X) &= \text{Median} \left(\min_{y \in Y} \delta(x, y) \right) \\ \rho_{Med}(X; Y) &= \text{Median} \left(\min_{x \in X} \delta(x, y) \right).\end{aligned}$$

Notice that the mean nearest neighbor distance is not a robust distance measure. That is, if for some reason a data point is far from the norm, the p-value computation becomes very sensitive to this data point. This can occur, for example, when a character in the real data set X is actually a ‘c’ (instead of being an ‘e,’), and is identified incorrectly as an ‘e’. Yet another outlier source is connected characters: when characters are extracted from a real document, they may touch other characters, pieces of which might slip in. The median and the trimmed mean distance measures are robust against outliers since

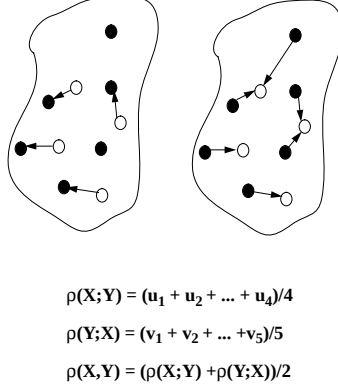


Figure 6: The black dots are elements of the set X and the white dots are elements of the set Y . In the figure on the left, the distance $\rho(X; Y)$ from Y to X is computed by summing the distance of each y_i to the nearest x_i . Similarly on the right the distance $\rho(Y; X)$ is computed. The final symmetric distance $\rho(X, Y)$ is computed by taking the mean.

they do not look at the tails of the distribution. One would expect that these measures should work better in cases where there are outliers.

The distance function $\delta(x, y)$ mentioned earlier is the distance between two individual characters x and y . We use the Hamming distance for δ . This is computed by counting the number of pixels where the characters x and y differ after the centroids of x and y have been registered. A variety of other character distances $\delta(x, y)$ and set distance functions $\rho(X, Y)$ could have been used (e.g. the Hausdorff distance, rank-ordered Hausdorff distance, etc.). The combination of character distance $\delta(x, y)$ and set distance $\rho(X, Y)$ that gives rise to the best power function is the best pair of character and set distances to use for the validation procedure.

5 Null Distribution for Gaussian Populations

In this section we compute the null distributions of two set distances $\rho(X, Y)$ when x_i and y_i are Gaussian distributed. We show that when they are each Gaussian distributed with a known variance σ^2 , the two distance functions considered are χ^2 distributed under the null hypothesis. Such closed-form solutions for the null distributions are possible only when the underlying distributions are known *a priori*. However, this is not the case in general — the Gaussian assumptions might be appropriate in some settings but completely wrong in other settings. Thus the non-parametric permutation method described in Section 4 is a much better approach to computing the null distributions when the forms of the sample distributions are not known. Nevertheless, for the purpose of validating the software and algorithm for computing the empirical null distribution, the Gaussian case is very useful since it allows us to compare the empirical distributions against known (theoretically computed) distributions.

5.1 Inter-Cluster Mean Distance

Let $X = \{x_1, x_2, \dots, x_N\}$ be a set such that $x_i \in R$ and $x_i \sim N(\mu_X, \sigma^2)$. Similarly, let $Y = \{y_1, y_2, \dots, y_N\}$ be a set such that $y_i \in R$ and $y_i \sim N(\mu_Y, \sigma^2)$. The problem is to test the null hypothesis that $\mu_X = \mu_Y$ when σ^2 is known.

Now we know that

$$\hat{\mu}_X = \frac{1}{N} \sum_{i=1}^N x_i \sim N(\mu_X, \sigma^2/N) \quad (5)$$

$$\hat{\mu}_Y = \frac{1}{N} \sum_{i=1}^N y_i \sim N(\mu_Y, \sigma^2/N) . \quad (6)$$

Therefore

$$\hat{\mu}_X - \hat{\mu}_Y \sim N(\mu_X - \mu_Y, 2\sigma^2/N) \quad (7)$$

and

$$\sqrt{N/2}(\hat{\mu}_X - \hat{\mu}_Y)/\sigma \sim N(\mu_X - \mu_Y, 1) . \quad (8)$$

Now let

$$t = \rho(X, Y) = \frac{N}{2\sigma^2}(\hat{\mu}_X - \hat{\mu}_Y)^2 .$$

Thus under the null hypothesis that $\mu_X = \mu_Y$, we have

$$t = \rho(X, Y) \sim \chi_1^2 . \quad (9)$$

Thus, instead of empirically computing the distributions as described in Section 4 we can use the above analytic form of the distribution to accept or reject the null hypothesis. Moreover, we see that the empirical method has reduced to a standard statistical technique when the underlying distribution is known to be Gaussian.

5.2 Likelihood Distance

In the previous section we picked a particular distance function $\rho(X, Y)$ and showed that its null distribution is χ_1^2 . In this section we pick a distance function based on the likelihood function of the data. It turns out that this distance function is the same as the one used in the previous section.

Let $X = \{x_1, x_2, \dots, x_N\}$, where $x_i \in R$, and $x_i \sim N(\mu_X, \sigma^2)$. Similarly, let $Y = \{y_1, y_2, \dots, y_N\}$, where $y_i \in R$, and $y_i \sim N(\mu_Y, \sigma^2)$. The problem is to test the null hypothesis that $\mu_X = \mu_Y = \mu$.

Let $\rho_Y(X)$ denote the distance of set X from set Y . Here we use a function of the likelihood for ρ :

$$\rho_X(Y) = f(P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma)) \quad (10)$$

$$\rho_Y(X) = f(P(x_1, \dots, x_N | y_1, \dots, y_N, \sigma)) . \quad (11)$$

In general, the above distances need not be symmetric in X and Y . Hence we also consider symmetric distances of the form

$$\rho(X, Y) = f(P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma)P(x_1, \dots, x_N | y_1, \dots, y_N, \sigma)) . \quad (12)$$

We can also consider the right-hand side in the equation above divided by $\log \max_{\mu} P(x_1, \dots, x_N, y_1, \dots, y_N | \mu, \sigma)$. That is,

$$\rho(X, Y) = \log \left(\frac{P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma) P(x_1, \dots, x_N | y_1, \dots, y_N, \sigma)}{\max_{\mu} P(x_1, \dots, x_N, y_1, \dots, y_N | \mu, \sigma)} \right). \quad (13)$$

We can use the standard rules of probability theory to manipulate the above equation as follows:

$$\begin{aligned} & P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma) \\ &= \int_{-\infty}^{\infty} P(\mu, y_1, \dots, y_N | x_1, \dots, x_N, \sigma) d\mu \\ &= \int_{-\infty}^{\infty} \frac{P(y_1, \dots, y_N, x_1, \dots, x_N, \mu, \sigma)}{P(x_1, \dots, x_N, \sigma)} d\mu \\ &= \int_{-\infty}^{\infty} \frac{P(y_1, \dots, y_N | x_1, \dots, x_N, \mu, \sigma) P(x_1, \dots, x_N, \mu, \sigma)}{P(x_1, \dots, x_N, \sigma)} d\mu \\ &= \int_{-\infty}^{\infty} \frac{P(y_1, \dots, y_N | \mu, \sigma) P(x_1, \dots, x_N | \mu, \sigma) P(\mu, \sigma)}{\int_{-\infty}^{\infty} P(x_1, \dots, x_N | \lambda, \sigma) P(\lambda, \sigma) d\lambda} d\mu. \end{aligned} \quad (14)$$

Now we make the assumption that μ and σ are independent so that $P(\mu, \sigma) = P(\mu)P(\sigma)$. Furthermore we assume that μ and σ have a uniform prior. Although this implies the prior is improper (since its integral is not equal to 1), the posterior distribution integrates to 1. Thus $P(\mu, \sigma) = P(\mu)P(\sigma) \propto \epsilon$. But the ϵ in the numerator and the denominator of equation (14) cancel out and the numerator can now be written as follows:

$$\begin{aligned} & P(y_1, \dots, y_N | \mu, \sigma) P(x_1, \dots, x_N | \mu, \sigma) P(\mu, \sigma) \\ &= \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^N e^{-\frac{1}{2\sigma^2} \sum_{i=1}^N (y_i - \mu)^2} \cdot \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^N e^{-\frac{1}{2\sigma^2} \sum_{j=1}^N (x_j - \mu)^2} \\ &= \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{2N} e^{-\frac{1}{2\sigma^2} \left[\sum_{i=1}^N (y_i - \mu)^2 + \sum_{j=1}^N (x_j - \mu)^2 \right]}. \end{aligned} \quad (15)$$

Since the denominator is not a function of either μ or $y_1, \dots, y_N$, it is a constant. The denominator can be computed by integrating out $\mu, y_1, \dots, y_N$ from the probability density in equation (15). Thus

$$\begin{aligned} & P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma) \\ &= C \int_{-\infty}^{\infty} \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{2N} e^{-\frac{1}{2\sigma^2} \left[\sum_{i=1}^N (y_i - \mu)^2 + \sum_{j=1}^N (x_j - \mu)^2 \right]} d\mu \end{aligned} \quad (16)$$

where the constant of integration C can be found by equating the right-hand side to 1. In order to compute the integral, we simplify the exponent inside the integral:

$$\sum_{i=1}^N (y_i - \mu)^2 + \sum_{j=1}^N (x_j - \mu)^2$$

$$\begin{aligned}
&= \sum_{i=1}^N (y_i - \bar{y} + \bar{y} - \mu)^2 + \sum_{j=1}^N (x_j - \bar{x} + \bar{x} - \mu)^2 \\
&= \sum_{i=1}^N (y_i - \bar{y})^2 + \sum_{i=1}^N (y_i - \bar{y})(\bar{y} - \mu) + N(\bar{y} - \mu)^2 \\
&\quad \sum_{j=1}^N (x_j - \bar{x})^2 + \sum_{j=1}^N (x_j - \bar{x})(\bar{x} - \mu) + N(\bar{x} - \mu)^2 \\
&= \sum_{i=1}^N (y_i - \bar{y})^2 + \sum_{j=1}^N (x_j - \bar{x})^2 + N(\bar{y} - \mu)^2 + N(\bar{x} - \mu)^2 .
\end{aligned} \tag{17}$$

But

$$\begin{aligned}
(\bar{y} - \mu)^2 + (\bar{x} - \mu)^2 &= \bar{x}^2 + \bar{y}^2 + 2 \left[\mu^2 - 2\mu \left(\frac{\bar{x} + \bar{y}}{2} \right) \right] \\
&= \bar{x}^2 + \bar{y}^2 - 2\mu \left(\frac{\bar{x} + \bar{y}}{2} \right)^2 \\
&\quad + 2 \left[\mu^2 - 2\mu \left(\frac{\bar{x} + \bar{y}}{2} \right) + \left(\frac{\bar{x} + \bar{y}}{2} \right)^2 \right] \\
&= \frac{(\bar{x}^2 + \bar{y}^2 - 2\bar{x}\bar{y})}{2} + 2 \left[\mu - \left(\frac{\bar{x} + \bar{y}}{2} \right) \right]^2 \\
&= \frac{(\bar{x} - \bar{y})^2}{2} + 2 \left[\mu - \left(\frac{\bar{x} + \bar{y}}{2} \right) \right]^2 .
\end{aligned} \tag{18}$$

Thus from equations (18) and (17)

$$\begin{aligned}
&\sum_{i=1}^N (y_i - \mu)^2 + \sum_{j=1}^N (x_j - \mu)^2 \\
&= \sum_{i=1}^N (x_i - \bar{x})^2 + \sum_{j=1}^N (y_j - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2 + 2N \left(\mu - \left(\frac{\bar{x} + \bar{y}}{2} \right) \right)^2 .
\end{aligned} \tag{19}$$

Also, since

$$\int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}(\sigma/\sqrt{2N})} e^{-\frac{1}{2\sigma^2/2N} \left(\mu - \frac{\bar{x} + \bar{y}}{2} \right)^2} d\mu = 1, \tag{20}$$

we have

$$\begin{aligned}
&P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma) \\
&= C \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{2N} \frac{\sqrt{2\pi}\sigma}{\sqrt{2N}} \cdot e^{\frac{1}{2\sigma^2/2N} \left[\sum_{i=1}^N (x_i - \bar{x})^2 + \sum_{j=1}^N (y_j - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2 \right]} .
\end{aligned} \tag{21}$$

Now to get the value of C we proceed as follows:

$$1 = C \int_{-\infty}^{\infty} \dots \int_{-\infty}^{\infty} \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{2N} e^{-\frac{1}{2\sigma^2} \left[\sum_{i=1}^N (y_i - \mu)^2 + \sum_{j=1}^N (x_j - \mu)^2 \right]} dy_1 \dots dy_N d\mu$$

$$\begin{aligned}
&= C \int \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^N e^{-\frac{1}{2\sigma^2} [\sum_{i=1}^N N(x_i - \bar{x})^2 + N(\mu - \bar{x})^2]} d\mu \\
&= C \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^N \frac{\sqrt{2\pi}\sigma}{\sqrt{N}} e^{-\frac{1}{2\sigma^2/N} [\sum_{i=1}^N (x_i - \bar{x})^2]}.
\end{aligned} \tag{22}$$

Thus we have computed C to be

$$C = \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{-(N+1)} \sqrt{N} e^{\frac{1}{2\sigma^2/N} [\sum_{i=1}^N (x_i - \bar{x})^2]}. \tag{23}$$

We now can write the complete conditional density as

$$\begin{aligned}
&P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma) \\
&= \left[\left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{-(N+1)} \sqrt{N} e^{\frac{1}{2\sigma^2/N} [\sum_{i=1}^N (x_i - \bar{x})^2]} \right] \\
&\quad \cdot \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^{2N} \frac{\sqrt{2\pi}\sigma}{\sqrt{2N}} \cdot e^{\frac{1}{2\sigma^2/2N} [\sum_{i=1}^N (x_i - \bar{x})^2 + \sum_{j=1}^N (y_j - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2]} \\
&= \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^N \cdot \sqrt{2} \cdot e^{-\frac{1}{2\sigma^2} [\sum_{i=1}^N (y_i - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2]}.
\end{aligned} \tag{24}$$

Thus we can use $2\sigma^2$ times the negative exponent of the conditional probability, as given in equation (24), as the test statistic $\rho_X(Y)$. Notice that it is not symmetric in X and Y .

$$\rho_X(Y) = f(P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma)) \tag{25}$$

$$= -\log P(y_1, \dots, y_N | x_1, \dots, x_N, \sigma) + \frac{N}{2} \log(2\pi\sigma^2) - \frac{1}{2} \log(2) \tag{26}$$

$$= \sum_{i=1}^N (y_i - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2. \tag{27}$$

$$\rho_Y(X) = f(P(x_1, \dots, x_N | y_1, \dots, y_N, \sigma)) \tag{28}$$

$$= -\log P(x_1, \dots, x_N | y_1, \dots, y_N, \sigma) + \frac{N}{2} \log(2\pi\sigma^2) - \frac{1}{2} \log(2) \tag{29}$$

$$= \sum_{i=1}^N (x_i - \bar{x})^2 + \frac{N}{2} (\bar{y} - \bar{x})^2. \tag{30}$$

In order to get a symmetric test statistic, we can look at the product of the conditional probabilities, so that

$$\rho_X(Y) + \rho_Y(X) = \sum_{i=1}^N (y_i - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2 + \sum_{i=1}^N (x_i - \bar{x})^2 + \frac{N}{2} (\bar{y} - \bar{x})^2. \tag{31}$$

But we know that the sum of the within-cluster scatter and the between-cluster scatter is equal to the total scatter. Thus

$$\sum_{i=1}^N (y_i - \bar{y})^2 + \frac{N}{2} (\bar{x} - \bar{y})^2 + \sum_{i=1}^N (x_i - \bar{x})^2 = \sum_{i=1}^N \left(x_i - \left(\frac{\bar{x} + \bar{y}}{2} \right) \right)^2 + \sum_{j=1}^N \left(y_j - \left(\frac{\bar{x} + \bar{y}}{2} \right) \right)^2.$$

Notice that for given data sets, the above summation is the same constant regardless of which points go with x_i and which with y_i . Thus

$$\rho_X(Y) + \rho_Y(X) = C + \frac{N}{2}(\bar{y} - \bar{x})^2 \quad (32)$$

where C is a constant. Thus a symmetric test statistic based on likelihood is

$$\rho(X, Y) = \frac{N}{2\sigma^2}(\bar{y} - \bar{x})^2. \quad (33)$$

The reason for normalizing by σ^2 will become clear shortly.

Monte Carlo hypothesis tests can now be conducted with the distance functions ρ defined in this section. In Figure 7 we show that the theoretically computed null distribution agrees with the null distribution computed empirically by random permutations.

It is important to statistically compare the test statistics $\rho_X(Y)$, $\rho_Y(X)$, and $\rho(X, Y)$ computed in this section. Notice that

$$\begin{aligned} \bar{x} &\sim N(0, \sigma^2/N), \\ \bar{y} &\sim N(0, \sigma^2/N), \\ \bar{x} - \bar{y} &\sim N(0, 2\sigma^2/N). \end{aligned}$$

Thus

$$(\bar{x} - \bar{y})^2 \sim \frac{2\sigma^2}{N} \chi_1^2$$

and

$$\rho(X, Y) = \frac{N}{2\sigma^2}(\bar{x} - \bar{y})^2 \sim \chi_1^2. \quad (34)$$

Thus $\rho(X, Y)$ has a mean of 1 and variance of 2. Similarly,

$$\frac{1}{\sigma^2} \sum_{i=1}^N (y_i - \bar{y})^2 \sim \chi_{N-1}^2$$

Thus

$$\frac{1}{\sigma^2} \sum_{i=1}^N (y_i - \bar{y})^2 + \frac{N}{2\sigma^2}(\bar{x} - \bar{y})^2 \sim \chi_{N-1}^2 + \chi_1^2,$$

so that

$$\rho_X(Y) \sim \chi_N^2. \quad (35)$$

We see that $\rho_X(Y)$ has a mean of N and variance of $2N$. This implies that $\rho(X, Y)$ is a more powerful test statistic (in terms of false alarms) than $\rho_X(Y)$ or $\rho_Y(X)$.

6 Experimental Protocol and Results

In this section we outline the protocol we use to conduct the experiments. Here we give all the sample sizes we use, the numbers of trials that are run at different stages, the exact model parameter values that are used for generating the synthetically degraded characters, the impact of the distance functions, etc. Three types of experiments are possible:

Synthetic vs. Synthetic: One sample X is synthetically created using the document degradation model, with a fixed model parameter value. Then many samples Y are generated, again using the model, but with different parameter settings. The validation procedure can be run on the samples X and Y , and the power function generated. This experiment is in part a sanity check for the methodology: if it does not work on controlled synthetic data, there is little point in trying it on real data.

Real vs. Real: This experiment tests for systematic dissimilarities between two image populations (e.g. rotations, fonts, etc.). Note that this use of the validation procedure is independent of the degradation model.

Real vs. Synthetic: Here the sample X consists of real degraded characters and the sample Y is generated by varying the degradation model parameter Θ . The validation procedure is run on the X and Y samples, and a power function is generated. This experiment tests whether or not the synthetic characters are actually close to the real characters.

6.1 Protocol for Synthetic vs. Synthetic

The following protocol is used for creating the samples X and Y . The distribution parameter Θ_X is fixed with the following parameter component values: $\eta_f = \eta_b = 0$, $\alpha_0 = \beta_0 = 1$, $\alpha = \beta = 1.5$, and structuring element size $k = 5$. The distribution parameter Θ_Y is varied by varying α and β . In our experiments we make α equal to β . The other parameter components of Θ_Y — $\eta_f, \eta_b, \alpha_0, \beta_0, k$ — are made equal to the corresponding components of the model parameter Θ_X . In all cases the noise-free document is the same (a L^AT_EX document page formatted in IEEE Transactions style) and the same set of 340 characters ‘e’ (Computer Modern Roman 10 point font) are extracted from the page to create the samples X and Y .

The validation procedure parameters used are as follows:

1. Sizes of samples X and Y : $N = M = \{10, 20, 60\}$.
2. Number of permutations: $K = 1000$.
3. Significance level of the test: $\epsilon = 0.05$.
4. Number of repetitions used in computing the power function: $T = 100$.
5. The character-to-character distance $\delta(x, y)$ used is the Hamming distance.
6. The set-to-set distance $\rho(X, Y)$ used is the mean nearest-neighbor distance.

The noise-free document is shown in Figure 8(a). The degraded document generated with model parameter Θ_X is shown in Figure 8(b). The power functions for sample sizes 10, 20, 60 are shown in Figure 9. The power function corresponding to sample size 10 is the widest, and the power function corresponding to sample size 60 is the narrowest. Note that all three power functions give a misdetection (reject) rate close to $\epsilon = 0.05$ when Θ_Y is close to Θ_X . (Only the α component, which is equal to 1.5 for Θ_X , is shown in the plot.)

Furthermore, when the α component for Θ_Y is far from 1.5, the misdetection rate is close to 1.0, which implies that the validation procedure can distinguish the two samples with high probability. An image generated with $\alpha = \beta = 1.7$ that the validation procedure accepted with a probability close to 0.9 is shown in Figure 8(c). Two document images generated with parameter values $\alpha = \beta = 2.0$ and $\alpha = \beta = 0.9$ that are easily rejected by the validation procedure are shown in Figure 8(d) and Figure 8(e), respectively.

6.2 Protocol for Real vs. Real Experiment

First, various European language texts are generated using the Adobe Times-Roman typeface at 8 point. Next, these documents are printed on a Canon laser printer and then scanned at 400 pixels per inch using a Canon scanner. Lower-case ‘e’s are extracted semiautomatically by OCR (thus some characters possess artifacts resulting from resegmentation). From among these, 3000 characters are selected by two persons working independently to avoid misclassifications.

Before selecting the two populations, we randomly shuffle the real data in order to obscure any systematic per-page dissimilarities (due to, for example, skew scale variations). The validation procedure does not reject the null hypothesis that the two samples are from the same underlying population. Repeated trials give a reject rate close to 0.05, the significance level designed into the test.

6.3 Outliers and Distance Function Comparisons

The validation procedure protocol is as follows: the significance level ϵ is fixed at 0.05; the sample sizes $N = M$ used are 10, 20, and 60; the number of permutations K for creating the empirical null distribution is 1000; the number of trials T for estimating the misdetection rate is 100.

We studied the sensitivity of the validation procedure to the set distance $\rho(X, Y)$ as follows. The data sets X and Y are collections of (synthetic) degraded characters ‘e’. The degradation parameter values for X are fixed at $\alpha = \beta = 1.5$, but the corresponding degradation parameters for Y are varied from 0.6 to 2.4. The Hamming distance is used for the character-to-character distance $\delta(x, y)$. The sample size of X and Y is fixed at $N = M = 60$. The mean, trimmed mean and median distances are used to compute the power function, in both the presence and absence of outliers.

Figures 10(a), 11(a), and 12(a) show the power functions in the absence of outliers when the mean and trimmed mean distances are used. Next, we introduced outliers into the data set X by replacing five degraded ‘e’s with degraded ‘c’s. The Y data set is unchanged. Figures 10(b), 11(b), and 12(b), show the power functions in the presence of outliers. Clearly the median and trimmed mean nearest neighbor distances are more robust against outliers, since the corresponding power functions are not affected. Furthermore, it can be seen that the median NN distance function, in the outlier-free case, is less “powerful” than the mean NN distance function since the median distance power function plot lies below the mean distance power function plot. Finally, it can be seen that the 10% trimmed NN distance function is superior to the other two distance functions, since the corresponding power function is robust against outliers and at the

same time higher.

6.4 Protocol for Validating Real vs. Synthetic Degradations

The ideal data is a $\text{\LaTeX}$ formatted document. The IEEE Transactions style is used for typesetting the document. The corresponding ideal binary image and character groundtruth are created using the DVI2TIFF software. The ideal document is created at 300×300 dots/inch resolution; the size of the binary document in pixels is 3300×2550 . This document is printed using a SparcPrinter II. Next, the original printed document is photocopied five times using a Xerox photocopier — once at the normal setting, twice with darker settings, and twice with lighter settings. Finally the five photocopied documents are scanned using a Ricoh scanner. The scanner is set at 300×300 dots/inch resolution. The rest of the scanner parameters are set at normal settings. The scanned binary image is of size 3307×2544 . The parameters are then estimated using the protocol specified in [12]. In all cases the noise-free document is the same (a $\text{\LaTeX}$ document page formatted in IEEE Transactions style) and the same set of 340 characters ‘e’ (Computer Modern Roman 10 point font) is extracted from the page to create the synthetic population Y .

The validation procedure parameters used are as follows:

1. Sample sizes of scanned characters X and synthetic characters Y : $N = M = \{10, 20, 60\}$.
2. Number of permutations for creating the empirical null distribution: $K = 1000$.
3. Significance level of the test: $\epsilon = 0.05$.
4. Number of bootstrap repetitions for computing the reject rate of the test: $T = 100$.
5. The bootstrap samples are sampled (with replacement) from a pool of size $N_b = 100$.
6. The character-to-character distance $\delta(x, y)$ used is the Hamming distance.
7. The set-to-set distance $\rho(X, Y)$ used is the mean nearest-neighbor distance.

The above test was conducted on ‘e’s. The test did not reject the null hypothesis that the samples are from the same population for a sample size of 10. That is, the reject rate is lower than 5%. For the sample size of 20, 46% of the time the test rejected the null hypothesis. For sample size of 60, the null hypothesis was rejected 100% of the time.

6.5 Comparing Two Models

In the previous section we used a two-sample permutation procedure to test the null hypothesis that the sample of real degraded characters and the sample generated by the estimated degradation model are from the same underlying population. We found that when the sample size is 40, the test procedure rejects the null hypothesis.

In fact, in a two-sample test, if one of the samples is from a distribution that is even slightly different from the second sample's distribution, the statistical testing procedure will be able to reject the null hypothesis that the samples are from the same underlying population if the sample size is large enough.

Since we know that any model of a real process, with very high likelihood, is an approximation to the real process, the samples generated from the model will be different from the real samples. Thus any validation procedure will be able to distinguish the real and synthetic samples if the sample sizes are large enough. In other words, it is futile to test the equality of the distributions of the synthetic samples and the real samples; they will always be proved to be unequal if a large enough sample size is used. Even if some other validation procedure is used, for example any method based on comparison of confusion matrices, the equality test is always going to give a negative result when the sample size is made large enough.

The next question is: How can one use the validation procedure in practice if the models are always going to be proved incorrect? The way to use the validation procedure is to compare two models and not evaluate just one model. That is, one can use the validation procedure to determine which model is closer to reality.

Suppose we have two document degradation models M_1 and M_2 . The problem is to find the model that is closer to the real process. We know that if the sample size N of the synthetic and real samples is increased, after a certain point the validation procedure will start rejecting both models. However, we will now give a procedure that will allow a researcher to decide which model is closer to reality for a fixed sample size N .

1. Fix the sample size N .
2. We are given a real sample D of size N .
3. Generate synthetic samples S_1 and S_2 of size N using the models M_1 and M_2 respectively.
4. Conduct the two-sample validation test using the real sample D and the synthetic sample S_1 . Let the associated p-value be p_1 .
5. Conduct the two-sample validation test using the real sample D and the synthetic sample S_2 . Let the associated p-value be p_2 .
6. If $p_1 > p_2$, model M_1 is closer to the real process for a sample size of N . Otherwise model M_2 is closer.

Thus the above procedure allows a researcher to choose between models. When we were choosing between parameter settings for a fixed model, we could use the power function to arrive at the best parameter sitting. However, two different models have different parameter spaces and hence they cannot be compared using power functions. The p-value provides a means of comparing the models on a common basis.

7 Conclusions

We have posed the degradation model validation problem as a statistical, two-sample, hypothesis testing problem. A non-parametric permutation test is used for this purpose. The user specifies a test statistic, which is essentially a distance function on the two sets of degraded characters. The null distribution of the test statistic, which is the distribution of the test statistic under the hypothesis that the two samples come from the same underlying population, is created using a permutation procedure. The p-value corresponding to the test statistic associated with the two sets is computed and compared with a user-specified significance level to reject or not reject the null hypothesis. This procedure and several robust variants are implemented and evaluated empirically. The goodness of the distance functions is evaluated using power functions, which are standard statistical devices. The local degradation model passes the validation test when the sample size is small but rejects it when the sample size is increased. This is so because any model of a real-world process is an approximation and thus will not pass the test if the sample size is increased. Another way of using the validation procedure is for choosing between models. After the validation procedure is run, a p-value is obtained. Thus if two different models are tested on the same real data, each validation procedure gives rise to a p-value for each model. The model whose associated p-value is larger is in closer agreement with the real data and thus should be preferred.

8 Acknowledgement

We would like to thank Azriel Rosenfeld of the University of Maryland for his comments. This research was funded in part by the Department of Defense and the Army Research Laboratory under Contract MDA 9049-6C-1250.

References

- [1] S. F. Arnold. *Mathematical Statistics*. Prentice-Hall, New Jersey, 1990.
- [2] H. Baird. Document image defect models. In *Proc. of IAPR Workshop on Syntactic and Structural Pattern Recognition*, pages 38–46, Murray Hill, NJ, June 1990.
- [3] H. Baird. Calibration of document image defect models. In *Proc. of Second Annual Symposium on Document Analysis and Information Retrieval*, pages 1–16, Las Vegas, NV, April 1993.
- [4] H. S. Baird. Document image defect models. In *Structured Document Image Analysis*. Springer-Verlag, New York, 1992.
- [5] G. Borgerfors. Distance transforms in digital images. *Computer Vision, Graphics, and Image Processing*, 34:344–371, 1986.
- [6] DARPA. *Proceedings of DARPA Workshop on Document Understanding*. Palo Alto, CA, 1992.

- [7] B. Efron and R. J. Tibshirani. *An Introduction to the Bootstrap*. Chapman and Hall, New York, 1993.
- [8] P. Good. *Permutation Tests: A Practical Guide to Resampling Methods for Testing Hypotheses*. Springer-Verlag, New York, 1994.
- [9] R. M. Haralick, I. Phillips, et al. UW-CDROM-I.
- [10] R.M. Haralick. Performance assessment of near perfect machines. *Machine Vision and Applications*, 2:1–16, 1989.
- [11] R.M. Haralick and L.G. Shapiro. *Computer and Robot Vision*. Addison-Wesley, Massachusetts, 1992.
- [12] T. Kanungo. *Document Degradation Models and a Methodology for Degradation Model Validation*. PhD thesis, University of Washington, Seattle, WA, 1996. <http://www.cfar.umd.edu/~kanungo/pubs/phdthesis.ps.Z>.
- [13] T. Kanungo and R. M. Haralick. Estimation of morphological degradation parameters. In *SPIE Proceedings*, volume 2424, pages 86–95, San Jose, CA, February 1995.
- [14] T. Kanungo, R. M. Haralick, H. S. Baird, W. Stuetzle, and D. Madigan. Document degradation models: Parameter estimation and model validation. In *Proc. of Int. Workshop on Machine Vision Applications*, Kawasaki, Japan, December 1994.
- [15] T. Kanungo, R. M. Haralick, and I. Phillips. Global and local document degradation models. In *Proc. of Second Int. Conf. on Document Analysis and Recognition*, pages 730–734, Tsukuba, Japan, October 1993.
- [16] T. Kanungo, R. M. Haralick, and I. Phillips. Non-linear local and global document degradation models. *Int. Journal of Imaging Systems and Technology*, 5:220–230, 1994.
- [17] T. Kanungo, M. Y. Jaisimha, J. Palmer, and R. M. Haralick. A quantitative methodology for analyzing the performance of detection algorithms. In *IEEE Int. Conf. on Computer Vision*, pages 247–252, Berlin, Germany, May 1993.
- [18] T. Kanungo, M. Y. Jaisimha, J. Palmer, and R. M. Haralick. A methodology for quantitative performance evaluation of detection algorithms. *IEEE Trans. on Image Processing*, 4:1667–1674, 1995.
- [19] D. E. Knuth. *TEX: the program*. Addison-Wesley, Massachusetts, 1988.
- [20] L. Lamport. *LATEX: a document preparation system*. Addison-Wesley, Massachusetts, 1986.
- [21] Y. Li, D. Lopresti, and A. Tomkins. Validation of document defect models for optical character recognition. In *Proc. of Third Annual Symposium on Document Analysis and Information Retrieval*, pages 137–150, Las Vegas, NV, April 1994.

- [22] Y. Li, D. Lopresti, and A. Tomkins. Validation of document defect models. *IEEE Trans. on Pattern Analysis and Machine Intelligence*, 18:99–107, 1996.
- [23] G. Nagy. Validation of OCR data sets. In *Proc. of Third Annual Symposium on Document Analysis and Information Retrieval*, pages 127–135, Las Vegas, NV, April 1994.
- [24] P. Vojta et al. XDVI Software, 1990.

Empirical and Theoretical Null Distribution

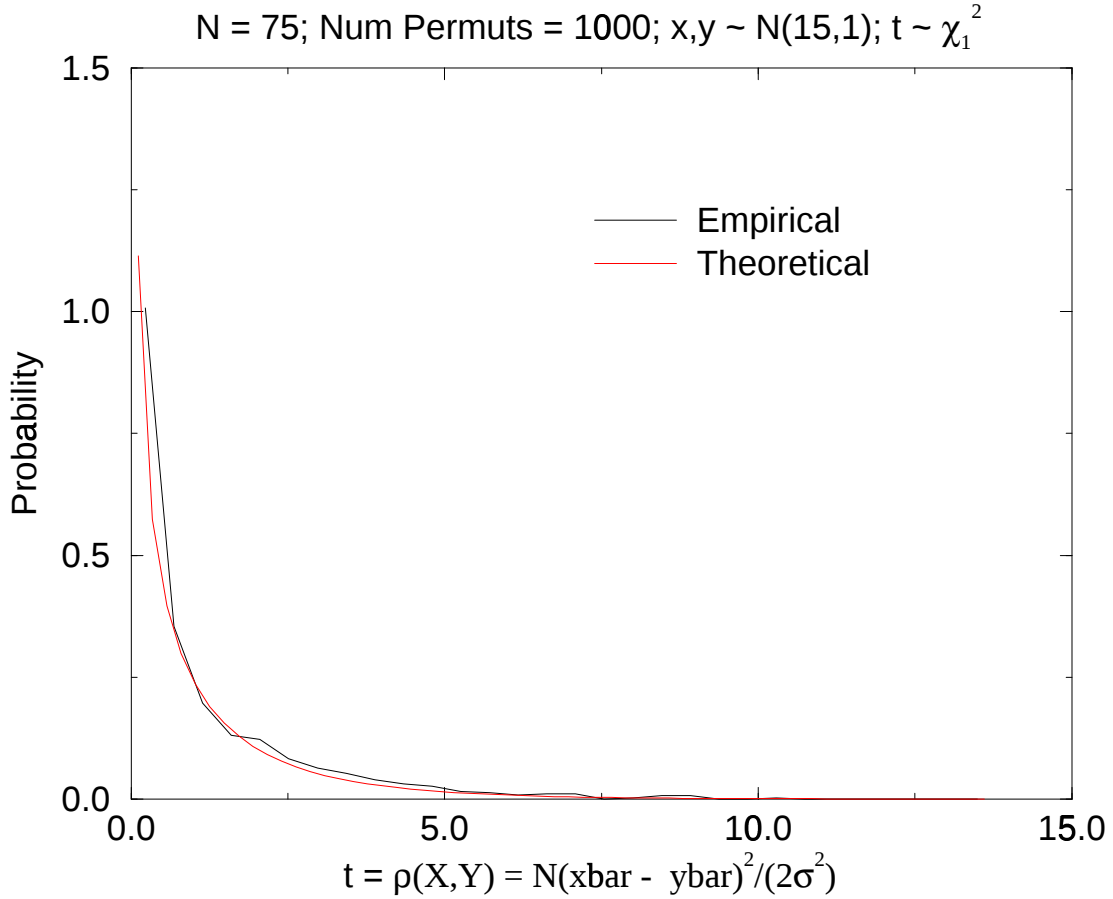


Figure 7: Empirical and theoretical null distributions for two sample tests. Samples X and Y of size $N = 75$ are drawn from $N(15, 1)$. The empirical null distribution is computed as described in Section 4. We use 1000 random permutations for computing the distribution. The distance function used is $t = \rho(X, Y) = N(\bar{x} - \bar{y})^2 / (2\sigma^2)$. The theoretical distribution of t is χ_1^2 . The empirical and theoretical plots have been plotted together in this figure.

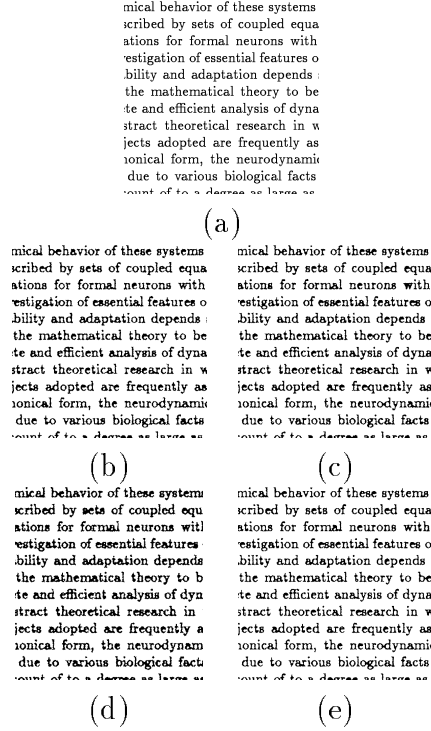


Figure 8: Local document degradation model. (a) Subimage of the noise-free document. (b) Degraded document generated with $\alpha = \beta = 1.5$. (c) A degraded image accepted as similar to (b), $\alpha = \beta = 1.7$; (d) A degraded image rejected as similar to (b), $\alpha = \beta = 0.9$; (e) A degraded image rejected as similar to (b), $\alpha = \beta = 2.0$. The sample size used is 60.

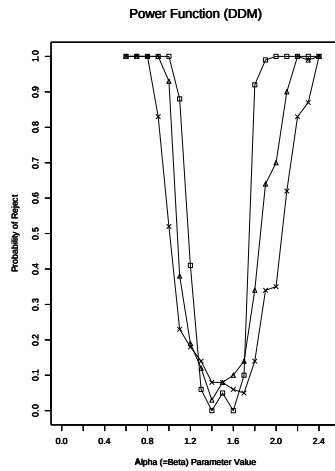


Figure 9: Power plots for the local document degradation model. The parameters for X were fixed with $\alpha = \beta = 1.5$, while the parameters for Y were varied. Notice that the power function has a minimum near $\alpha = \beta = 1.5$. The power function corresponding to a sample size of 60 (boxes) is sharper; that corresponding to a sample size of 10 (crosses) is broader.

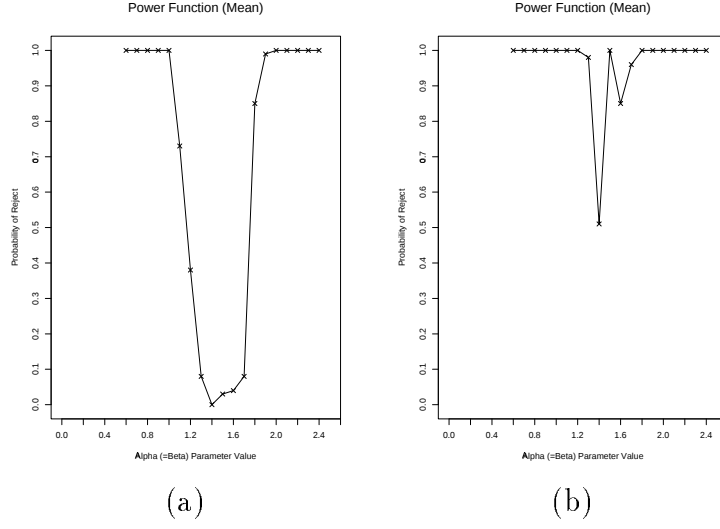


Figure 10: Power functions of the validation procedure when mean nearest neighbor distance is used for the set distance function $\rho(X, Y)$. Figure (a) is when there are no outliers. Figure (b) corresponds to the situation when there are five outliers in one of the data sets.

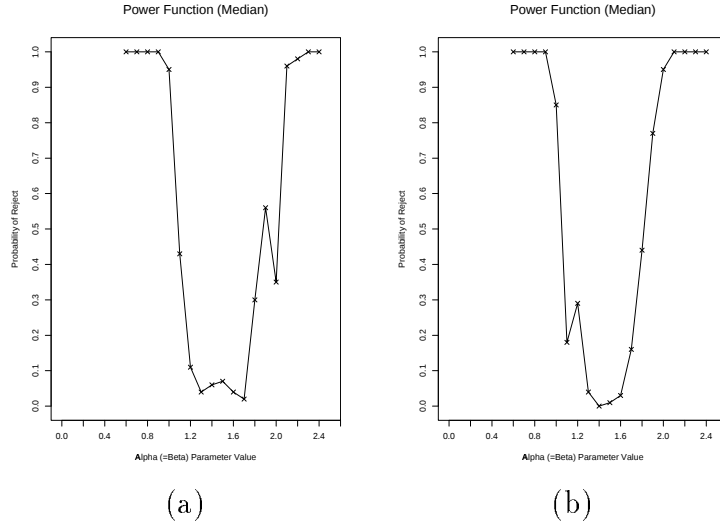


Figure 11: Power functions of the validation procedure when median nearest neighbor distance is used for the set distance function $\rho(X, Y)$. Figure (a) is when there are no outliers. Figure (b) corresponds to the situation when there are five outliers in the X data set.

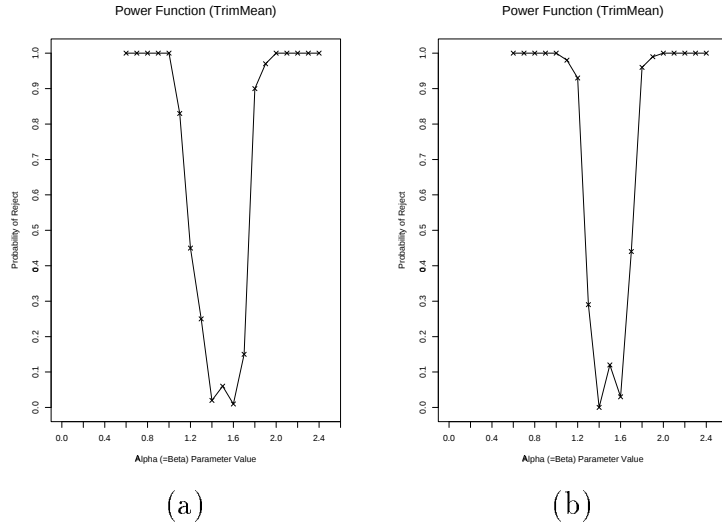


Figure 12: Power functions of the validation procedure when 10% trimmed mean nearest neighbor distance is used for the set distance function $\rho(X, Y)$. Figure (a) is when there are no outliers. Figure (b) corresponds to the situation when there are five outliers in the X data set.