

TEST THE EFFECT OF ENHANCEMENT FILTERS ON THE NOISY TEXT IMAGE

Nidhal K. shukur

Assistance lecturer, Computer and SW Engineering Department
Al- Mustansyria University

ABSTRACT

This paper explains an important feature of OCR (Optical Character Recognition) software and makes a comparison between them and shows how the enhancement filters effected on noisy text image, using a Hybrid Method (HM) with verifying factor with Gaussian filter gives a good result with BNTI (Black Noisy Text Image) removal, also with hybrid mean filter with verifying factor gives noise and blurring removal.

Also the results show that with EBNTI (Enhanced Black Noisy Text Image) the number of error is reduced than actually exit in the original BNTI and median filter work best with BNTI image than with WNTI (Wight Noisy Text Image) but still these filter causes some noise removal but do not enhance all the defected characters.

Keywords: Hybrid Method, BNTI, WNTI, EBNTI, OCR

Introduction

Optical character recognition, usually abbreviated to OCR, is the mechanical or electronic translation of scanned images of handwritten, typewritten or printed text into machine-encoded text. It is widely used to convert books and documents into electronic files, to computerize a record-keeping system in an office, or to publish the text on a website. OCR makes it possible to edit the text, search for a word or phrase, store it more compactly, display or print a copy free of scanning artifacts, and apply techniques such as machine translation, text-to-speech and text mining to it. OCR is a field of research in pattern recognition, artificial intelligence and computer vision.

OCR systems require calibration to read a specific font; early versions needed to be programmed with images of each character, and worked on one font at a time. "Intelligent" systems with a high degree of recognition accuracy for most fonts are now common. Some systems are capable of reproducing formatted output that closely approximates the

original scanned page including images, columns and other non-textual components. [1]

1.1 What is OCR?

OCR stands for optical character recognition. Optical character recognition technology makes it possible to automatically recognize text in scanned documents, so that users can work with such files as easily as they can with files created on a computer. OCR uses the outlines of letters and numbers to identify text content. While OCR is limited to printed text and cannot recognize handwriting, it is extremely effective for this category of text [2].

1.2 What determines the quality of OCR?

The primary determinant of OCR quality is accuracy. A good OCR program requires a high level of accuracy, so that users can be confident that they can effectively work with the text in a scanned document. Another important feature of OCR technology is volume capability. For industrial purposes, it is often necessary to process large volumes of documents, and so it is essential to be able to handle large numbers of files quickly for more detail see the previous reference. It is very important to digitize (historic) paper based documents such as books and newspapers, has emerged. The aim is to preserve these documents and make them fully accessible, searchable and process able in digital form. Knowledge contained in paper based documents is more valuable for today's digital world when it is available in digital form. The first step towards transforming a paper based archive into digital archives to scan the documents. The next step is to apply an OCR (Optical Character Recognition) process, meaning that the scanned image of each document will be translated into machine process able text. Due to the print quality of the documents and the error-prone pattern matching techniques of the OCR process, OCR errors occur. Modern OCR processors have character recognition rates up to 99% on high quality documents. Assuming an average word length of 5 characters, this still means that one out of 20 words is defect. Thus, at least 5% of all processed words will contain OCR errors. On historic documents the error rate will be even higher because the print quality is likely to be of lower quality [3].

1.3 Finding super high quality OCR

If you need a super high quality OCR solution, you need to choose a program designed for professional use. Some of these can be expensive, so it is important to be confident about the software that you want before purchasing it. Fortunately, CVISION offers a free trial version of Maestro

Recognition Server with OCR solution for more detail see the previous reference.

2. History

In 1929 Gustav Tauschek obtained a patent on OCR in Germany, followed by Handel who obtained a US patent on OCR in USA in 1933 (U.S. Patent 1,915,993). In 1935 Tauschek was also granted a US patent on his method (U.S. Patent 2,026,329). Tauschek's machine was a mechanical device that used templates and a photodetector. RCA engineers in 1949 worked on the first primitive computer-type OCR to help blind people for the US Veterans Administration, but instead of converting the printed characters to machine language, their device converted it to machine language and then spoke the letters. It proved far too expensive and was not pursued after testing for more detail see the previous reference.

In 1950, David H. Shepard, a cryptanalyst at the Armed Forces Security Agency in the United States, addressed the problem of converting printed messages into machine language for computer processing and built a machine to do this, reported in the Washington Daily News on 27 April 1951 and in the New York Times on 26 December 1953 after his U.S. Patent 2,663,758 was issued. Shepard then founded Intelligent Machines Research Corporation (IMR), which went on to deliver the world's first several OCR systems used in commercial operation. The first commercial system was installed at the Reader's Digest in 1955. The second system was sold to the Standard Oil Company for reading credit card imprints for billing purposes. Other systems sold by IMR during the late 1950s included a bill stub reader to the Ohio Bell Telephone Company and a page scanner to the United States Air Force for reading and transmitting by teletype typewritten messages. IBM and others were later licensed on Shepard's OCR patents. In about 1965 Reader's Digest and RCA collaborated to build an OCR Document reader designed to digitize the serial numbers on Reader's Digest coupons returned from advertisements. The font used on the documents were printed by an RCA Drum printer using the OCR-A font. The reader was connected directly to an RCA 301 computer (one of the first solid state computers). This reader was followed by a specialized document reader installed at TWA where the reader processed Airline Ticket stock. The readers processed document at a rate of 1,500 documents per minute, and checked each document, rejecting those it was not able to process correctly. The product became part of the RCA product line as a reader designed to process "Turn around Documents" such as those Utility and insurance bills returned with payments. The United States Postal Service has been using OCR machines to sort mail since 1965 based on

technology devised primarily by the prolific inventor Jacob Rabinow. The first use of OCR in Europe was by the British General Post Office (GPO). In 1965 it began planning an entire banking system, the National Giro, using OCR technology, a process that revolutionized bill payment systems in the UK. Canada Post has been using OCR systems since 1971[citation needed]. OCR systems read the name and address of the addressee at the first mechanized sorting center, and print a routing bar code on the envelope based on the postal code. To avoid confusion with the human-readable address field which can be located anywhere on the letter, special ink (orange in visible light) is used that is clearly visible under ultraviolet light. Envelopes may then be processed with equipment based on simple barcode readers. In 1974 Ray Kurzweil started the company Kurzweil Computer Products, Inc. and led development of the first omni-font optical character recognition system—a computer program capable of recognizing text printed in any normal font. He decided that the best application of this technology would be to create a reading machine for the blind, which would allow blind people to have a computer read text to them out loud. This device required the invention of two enabling technologies—the CCD flatbed scanner and the text-to-speech synthesizer. On January 13, 1976 the successful finished product was unveiled during a widely-reported news conference headed by Kurzweil and the leaders of the National Federation of the Blind. In 1978 Kurzweil Computer Products began selling a commercial version of the optical character recognition computer program. LexisNexis was one of the first customers, and bought the program to upload paper legal and news documents onto its nascent online databases. Two years later, Kurzweil sold his company to Xerox, which had an interest in further commercializing paper-to-computer text conversion. Kurzweil Computer Products became a subsidiary of Xerox known as Scansoft, now Nuance Communications [1].

3. OCR technology

The accurate recognition of Latin-script, typewritten text is now considered largely a solved problem on applications where clear imaging is available such as scanning of printed documents. Typical accuracy rates on these exceed 99%; total accuracy can only be achieved by human review. Other areas—including recognition of hand printing, cursive handwriting, and printed text in other scripts (especially those East Asian language characters which have many strokes for a single character)—are still the subject of active research.

Accuracy rates can be measured in several ways, and how they are measured can greatly affect the reported accuracy rate. For example, if word context (basically a lexicon of words) is not used to correct software finding non-existent words, a character error rate of 1% (99% accuracy) may result in an error rate of 5% (95% accuracy) or worse if the measurement is based on whether each whole word was recognized with no incorrect letters.[4]

On-line character recognition is sometimes confused with Optical Character Recognition[5] (see Handwriting recognition). OCR is an instance of off-line character recognition, where the system recognizes the fixed static shape of the character, while on-line character recognition instead recognizes the dynamic motion during handwriting. For example, on-line recognition, such as that used for gestures in the Penpoint OS or the Tablet PC can tell whether a horizontal mark was drawn right-to-left, or left-to-right. On-line character recognition is also referred to by other terms such as dynamic character recognition, real-time character recognition, and Intelligent Character Recognition or ICR. On-line systems for recognizing hand-printed text on the fly have become well-known as commercial products in recent years. Among these are the input devices for personal digital assistants such as those running Palm OS. The Apple Newton pioneered this product. The algorithms used in these devices take advantage of the fact that the order, speed, and direction of individual lines segments at input are known. Also, the user can be retrained to use only specific letter shapes. These methods cannot be used in software that scans paper documents, so accurate recognition of hand-printed documents is still largely an open problem. Accuracy rates of 80% to 90% on neat, clean hand-printed characters can be achieved, but that accuracy rate still translates to dozens of errors per page, making the technology useful only in very limited applications.

Recognition of cursive text is an active area of research, with recognition rates even lower than that of hand-printed text. Higher rates of recognition of general cursive script will likely not be possible without the use of contextual or grammatical information. For example, recognizing entire words from a dictionary is easier than trying to parse individual characters from script. Reading the Amount line of a cheque (which is always a written-out number) is an example where using a smaller dictionary can increase recognition rates greatly. Knowledge of the grammar of the language being scanned can also help determine if a word is likely to be a verb or a noun, for example, allowing greater accuracy. The shapes of individual cursive characters themselves simply do not contain

TEST THE EFFECT OF ENHANCEMENT FILTERS ON THE NOISY TEXT IMAGENidhal K. shukur

enough information to accurately (greater than 98%) recognize all handwritten cursive script. It is necessary to understand that OCR technology is a basic technology also used in advanced scanning applications. Due to this, an advanced scanning solution can be unique and patented and not easily copied despite being based on this basic OCR technology. For more complex recognition problems, intelligent character recognition systems are generally used, as artificial neural networks can be made indifferent to both affine and non-linear transformations [5] .

A technique which is having considerable success in recognizing difficult words and character groups within documents generally amenable to computer OCR is to submit them automatically to humans in the RECAPTCHA system. An OCR SDK is a Software Development Kit for adding Optical character recognition capabilities to forms processing applications, document imaging management systems, e-discovery systems and records management solutions. In order to avoid the difficulties of incorporating OCR technology, some OCR SDKs contain a high number of APIs, support multiple operating systems and programming languages [1].

Notes	Linux	Windows	Online	Name
ABBYY also supplies SDKs for embedded or mobile devices. Professional, Corporate and Site License Editions for Windows, Express Edition for Mac.	Yes	Yes	Yes	ABBYY FineReader
Works with structured, semi-structured, and unstructured documents.	No	Yes	No	AnyDoc Software
Template-free data extraction and processing of data from documents into any backend system; sample document types include invoices, remittance statements, bills of lading and POs.	No	Yes	No	Brainware
Enterprise-class system, can save text formatting and recognizes complicated tables of any structure	Yes	Yes	No	CuneiForm/OpenOCR
Won the highest marks in the independent testing performed by UNLV for X consecutive years (in 1994). The speed of ExperVision's OpenRTK is four to eight times faster than competition. — PC Magazine[citation needed] but also "Not as accurate as rival products, clumsy interface, limited options for proofreading, couldn't open some files in standard PDF or image formats."	Yes	Yes	Yes	ExperVision TypeReader & RTK
	Yes	No	No	GOOCR
Image to OCR Converter can also create searchable PDF and a unique text-only PDF file format when	No	Yes	No	Image to OCR Converter

TEST THE EFFECT OF ENHANCEMENT FILTERS ON THE NOISY TEXT IMAGENidhal K. shukur

used in conjunction with Universal Converter.				
	Yes	Yes	Yes	Img2txt.ru
Supports Latin, Asian, Arabic, and MICR character sets. For full page, zonal, and form image processing. Includes OCR, barcode, OMR and forms recognition. ICR (handwritten text recognition) is supported.	No	Yes	No	LEADTOOLS
Uses OmniPage	No	Yes	No	Microsoft Office Document Imaging
	No	Yes	No	Microsoft Office OneNote 2007
Command line	Yes	Yes	Yes	Ocrad
Pluggable framework which can use Tesseract	Yes	No	No	OCROPUS
Product of Nuance Communications	No	Yes	No	OmniPage
Automated data capture from structured, semi-structured, and unstructured documents. DocXP overview	No	Yes	No	Peladon Software DocXP
Capture structured and semi-structured documents from any scanner, MFP or network folder. Output to SharePoint and over 40 different backend systems PSI:Capture overview	No	Yes	No	PSIGEN Software PSI:Capture
.NET OCR SDK based on Cognitive Technologies' CuneiForm recognition engine. Wraps Puma COM server and provides simplified API for .NET applications	No	Yes	No	Puma.NET
Product of I.R.I.S. Group of Belgium. Asian and Middle Eastern editions.	No	Yes	No	Readiris
Scan, capture and classify business documents such as invoices, forms and purchase orders integrated with business processes.	No	Yes	No	ReadSoft
Converts faxed pages into editable document formats (doc, pdf, etc...).	No	Yes	No	RelayFax
For working with localized interfaces, corresponding language support is required.	No	Yes	No	Scantron
	No	Yes	No	SimpleOCR
For musical scores	No	Yes	No	SmartScore
Created by Hewlett-Packard; under further development by Google	Yes	Yes	No	Tesseract
	No	Yes	No	Transym OCR
	No	Yes	No	Zonal OCR

Tabel 1: a non-exhaustive comparison of optical character recognition software

4. Several Ocr software

OCR Software and ICR Software technology are analytical artificial intelligence systems that consider sequences of characters rather than whole words or phrases. Based on the analysis of sequential lines and

curves, OCR and ICR make 'best guesses' at characters using database look-up tables to closely associate or match the strings of characters that form words [1].

4.1. Free OCR Software

Webocr also known as Web-based OCR service or Online OCR service has been a new trend to meet larger volume and larger group of users after 30 years development of the desktop OCR. Internet and broad band technologies have made Webocr practically available to both individual users and enterprise customers. Free-OCR.com is a free online OCR (Optical Character Recognition) tool. You can use this service to extract text from any image you supply. This service is free, no registration necessary. We also do not need your email address. Just upload your image files. Free-OCR takes either a JPG, GIF, TIFF BMP or PDF (only first page). The only restriction is that the images must not be larger than 2MB, no wider or higher than 5000 pixels and there is a limit of 10 image uploads per hour. Free-OCR can handle images with multi-column text and also supports more languages: Bulgarian, Catalan, Czech, Danish, Dutch, English, Finnish, French, German, Greek, Hungarian, Indonesian, Italian, Latvian, Lithuanian, Norwegian, Polish, Portuguese, Romanian, Russian, Serbian, Slovak, Slovene, Spanish, Swedish, Tagalog, Turkish, Ukrainian, Vietnamese. Are you looking for ways to extract text and images from a handwritten, printed or typewritten document? The easiest way to do it is to scan the document and use an **Optical Character Recognition (OCR)** software to extract the content. There are several OCR softwares in the market. But many are commercial softwares and there are only a few freewares. We earlier covered Simple OCR, a free OCR software to read, and convert a hard copy document with standard fonts, into an editable soft copy document. Simple OCR is available as both a freeware and as a commercial version. Here we cover Free OCR, which is both a **Scanner Software** and an **OCR Software**. It is thus a complete scan and OCR program that includes the Windows compiled Tesseract free OCR engine, also known as a Tesseract GUI. Free OCR is not only free but is also very easy to use. Free OCR supports Optical Character Recognition (OCR) of multi-page Tiff, Adobe PDF and fax documents, as well as most image types including compressed Tiff. Free OCR for scanned PDF is based on Tesseract OCR PDF engine, an open source product released by Google. To use Free OCR, you should have .Net Framework 2.0 installed on your PC. The underlying Tesseract OCR engine requires images at a resolution of 200 dpi or greater and as such it is not suited for reading PC screenshots, which are only about 72dpi. The developers recommend scanning the

TEST THE EFFECT OF ENHANCEMENT FILTERS ON THE NOISY TEXT IMAGENidhal K. shukur

documents at 300 dpi grey scale (optimal level), for best results. Free OCR cannot read images that are upside down or rotated by 90 Degrees. Hence, make use of the rotate buttons to rotate the images before using Free OCR on them. You can also select the text area for Optical Character Recognition (OCR), by drawing a box around it. This gives better results than trying to OCR whole pages. Free OCR only supports scanned PDFs ie. PDF's that contain an image and gives better results from Clean scans. [1]

4.2. ABBYY FineReader Professional 10.501.159.70013

ABBYY FineReader Professional 10 is professional OCR (optical character recognition) software created to recognize text and to convert scanned paper documents, digital photographs and PDFs to editable and searchable electronic files. This software offers a superior recognition accuracy, it recreates the logical structure of documents, recognizes text columns, page numbering, headers and footers, footnotes, table of contents, font type and style, table structures, bookmarks, hyperlinks, vertical text, etc.

The application recognizes (auto-detects) text written in several languages (uses a built-in spell check), supports multilingual documents, is able to recognize low quality documents (faxes, text in photos), barcodes and converts PDF files to editable and searchable formats. After launching, users can scan their paper documents and photos with a built-in scanning tool, then save them to Word, Excel, searchable or editable PDF, image, HTML and other file formats. Texts from already scanned documents and images can be recognized and converted to DOC, PDF, HTML, XLS, DCX, DJVU, DJV, WDP, etc. formats.

ABBYY FineReader Professional has an easy to use interface, it is able to search and replace the recognized text, split and merge table cells, rotate, flip pages, email and print text and images.

The usage of the software is very simple, just select a task from the list, scan your paper documents or photos or open already scanned files, wait until the application recognizes the text and sends it to a word processor for editing or for viewing. Users can also see and manage the results in the right panel of the application, select a text, image, table or barcode area and launch the text recognition process again or edit the opened scanned document with a built-in image editor (image deskwing, correction, rotating, flipping, splitting, cropping, inverting, resolution changing, pixel erasing and distortion correction are supported).

An Automation Manager lets you set the application to automatically scan and convert documents. Other features included in this software are image previewing, text editing, multi-core processor support,

fast document reading and text recognition and automatic image processing (it splits dual pages, detects page orientation, corrects image resolution and 3D perspective distortion, straightens curved text lines).

Finally, this professional character recognition software is ideal both for beginners and professional, it recognizes text in scanned documents and photos, offers high quality outputs, performs several document conversion, creates and executes automated tasks, is available in multiple languages and is compatible with Windows 7. The demo version of ABBYY FineReader Professional 10 can be tested 15 days from the installation date, it can process up to 50 pages and can save up to 1 page at a time [6].

4.3. Music OCR

Music OCR is the application of optical character recognition to interpret sheet music or printed scores into editable and, often, playable form. Once captured digitally, the music can be saved in commonly used file formats, e.g. MIDI (for playback) and MusicXML (for page layout) [7].

4.3.1 History

Early research into recognition of printed sheet music was performed at the graduate level in the late 1960s at MIT and other institutions. Successive efforts were made to localize and remove musical staff lines leaving symbols to be recognized and parsed. The first commercial music-scanning product, MIDISCAN, was released in 1991 by Musitek corporation, for more detail, see the previous reference.

Unlike OCR of text, where words are parsed sequentially, music notation involves parallel elements, as when several voices are present along with unattached performance symbols positioned nearby. Therefore, the spatial relationship between notes, expression marks, dynamics, articulations and other annotations is an important part of the expression of the music. Modern music OCR packages have accuracy exceeding 99% when a clean scan is used and the notation is not exceptional (e.g. unfilled voices, non-standard symbology, etc.). Because music notation utilizes dots for staccato marks or to extend the value of a note, artifacts in the scan can lead to interpretation problems [8].

5. OCR Accuracy

Ocr accuracy levels must be very high to avoid unacceptable error correction workloads. An A4 page with 2000 marks recognized With 90 per Accuracy will contain 200 errors. At five second per error correction will require approximately 17, compared to the fact that a professional typist will write the entire page in less than 12 minutes. [9]

If the text is not grasp and clear, or there are stains and dirt, these are likely to cause trouble in scanning. E.g. Pages produced by dot matrix printers are likely to be problematic. Too much space between dots produces a broken character, which results in false recognition. Other material likely to cause trouble are old photo-copies, as copying smears text edges, and tinted and colored paper, as color reduces the contrast between background and text and makes it difficult to separate the text from the noise. The most cost-effective way is to retype it. This takes less time than correcting a text file full of errors. [10]

6. Experimental study

The nature of the degradation can affect the accuracy of the resultant OCR effort [10]. In this paper we concentrate on applying various enhancement filters on the noisy text image, the implemented algorithm as follow:

1. Apply a smoothing enhancement algorithms to produce an enhanced image for the noisy text image.
2. Take scan of for the text image .
3. Record the error characters.
4. Save the result text as image.
5. Compare the enhanced image with the original.
6. Repeat steps 1 to 5 by using a hybrid method and compare the new results with the old result and save the best.
7. End.

Figures 1-2 show sections from example pages of noise effects. Text image with white noise causes character break and text image with black noise causes character touching and text image with black and white noise causes the two previous trouble.

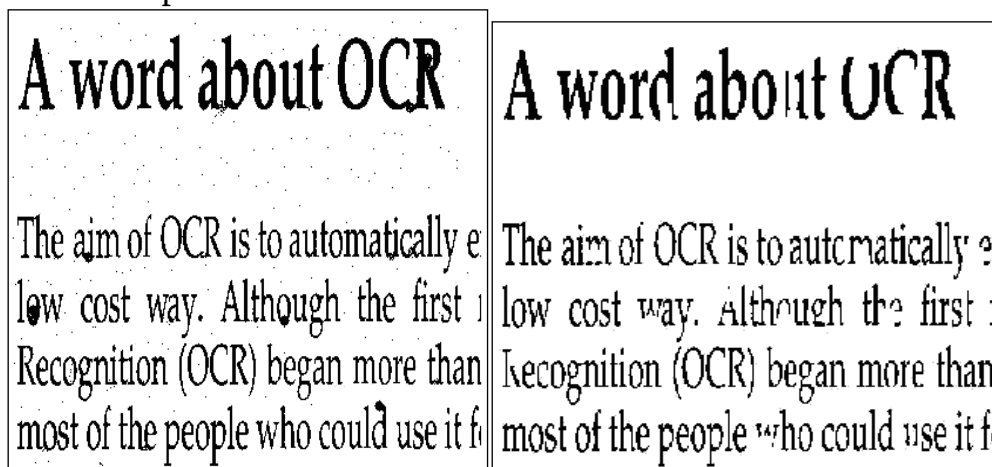


Fig. 1 text image with black noise

Fig.2 text image with white noise

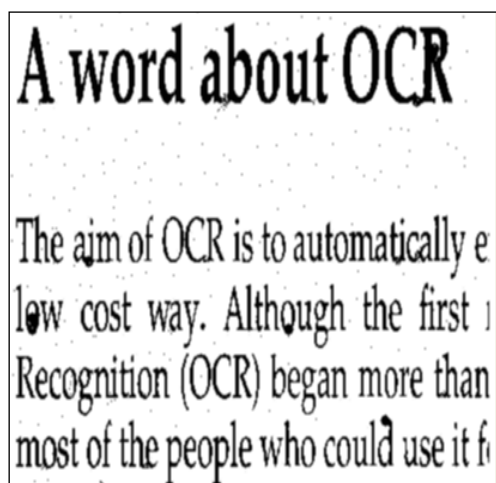


Fig. 3 enhanced text image with mean 3

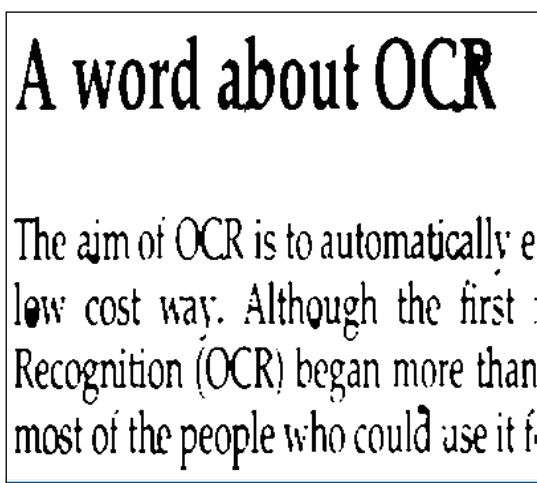


Fig.4 enhanced text image with median 3*3 filter

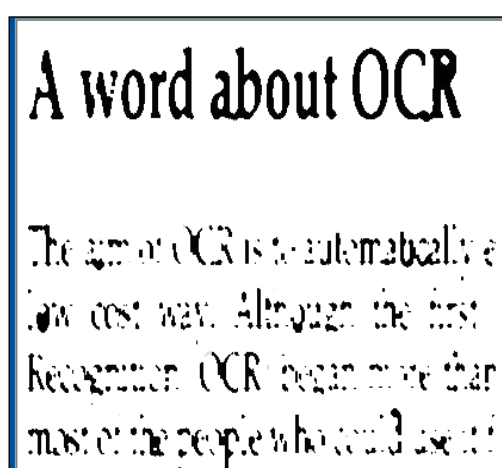


Fig. 5 enhanced text image with median 5*5 filter

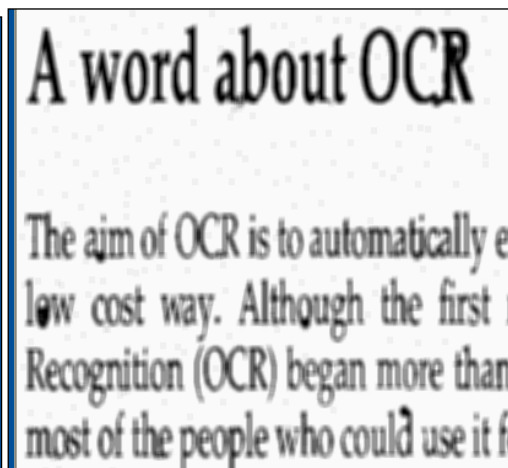


Fig. 6 enhanced text image with Gaussian filter

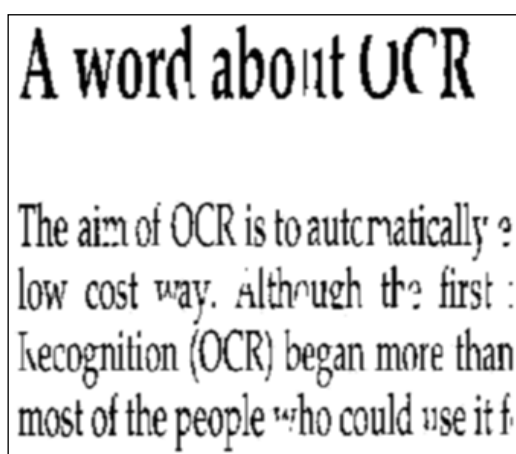


Fig. 7 enhanced text image with mean 3*3 filter

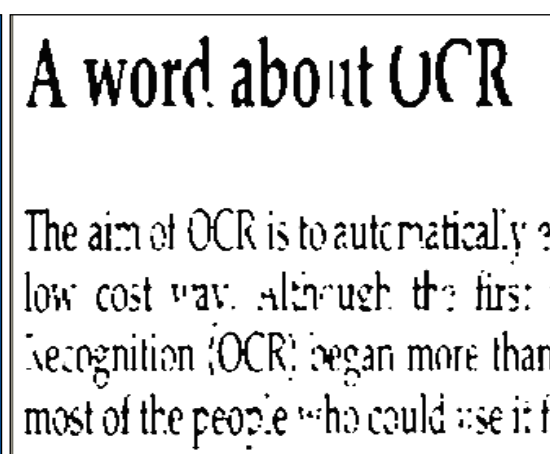


Fig. 8 enhanced text image with median 3*3 filter

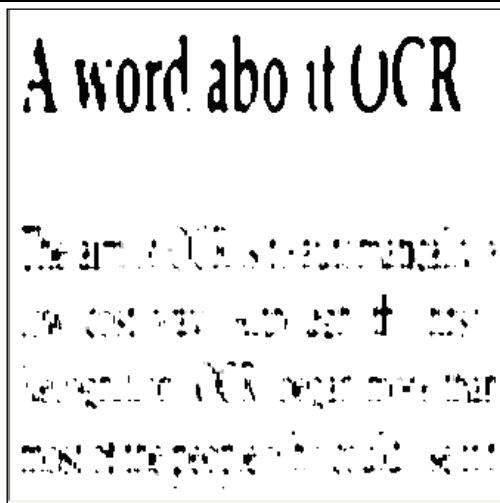


Fig. 9 enhanced text image with median 5*5 filter

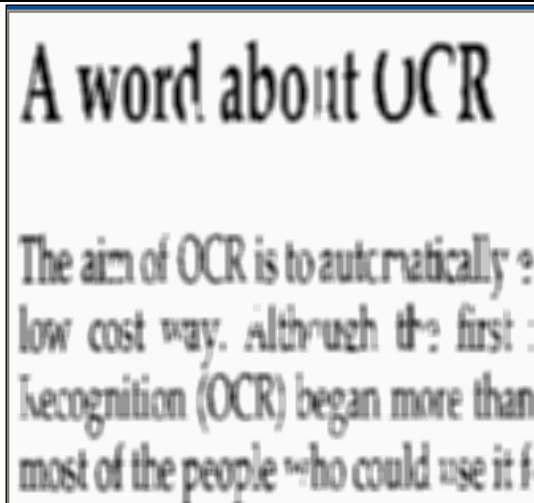


Fig. 10 enhanced text image with Gaussian filter

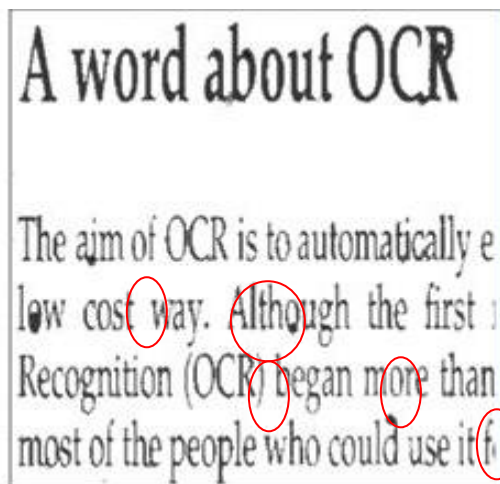


Fig. 11 a hybrid method using mean 3*3 filter

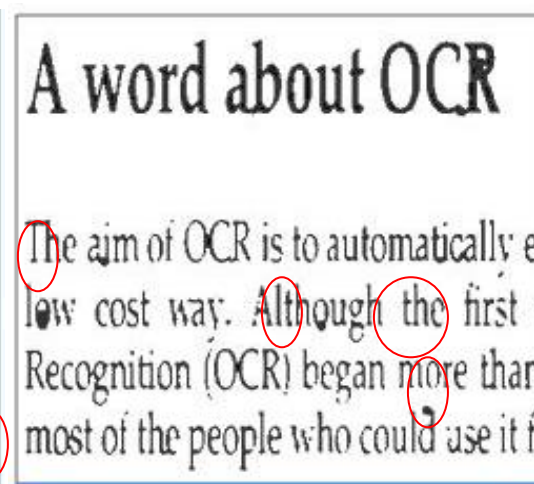


Fig. 12 a hybrid method using median 3*3 filter

Figures 3 to 10 show section from example page of noise effect enhanced by an enhancement algorithms such that (mean 3*3, mean 5*5, median 3*3, median 5*5, Gaussian). but OCR results are still poor on documents like these.

Figure 3-6 enhanced by mean 3*3, median 3*3, median 5*5, Gaussian methods but still there are characters touching problem. The OCR result shows there is only one error with "low" word in the character "o" recognized as "I" character so it is good result because there are 6 errors in the original BNTI, the same result with median filter, but with better cleaner image, so median filter show the best results, and with the Gaussian gives worse results with huge no. of OCR errors.

Fig. 11-12, illustrate the result after using a hybrid method (smoothing with edge detection methods) with verifying factor value, the experimental result shows that using a hybrid method with factor=.05 with gaussian filter gives a good result with BNTI removal and with only one error, also with

TEST THE EFFECT OF ENHANCEMENT FILTERS ON THE NOISY TEXT IMAGENidhal K. shukur

mean filter with factor=.2 gives noise and blurring removal and with only one error.

Fig. 6-10, enhanced by the above filters but still there are characters gaps problem and they are increased in addition to the previous founded gaps especially in median filter and OCR errors are widely increased with median 5*5 and Gaussian filters, so the previous character gaps are increased after enhancement like that in appeared in automatically , although, began and way wordsetc.

Seeing that how the characters gaps are decreased by a hybrid mean method compared with the a hybrid median filter as indicated in the fig.11 and fig.12 .

Chart 1 and chart 2 in Fig. 13 and Fig. 14 respectively , summarized all above explanation.

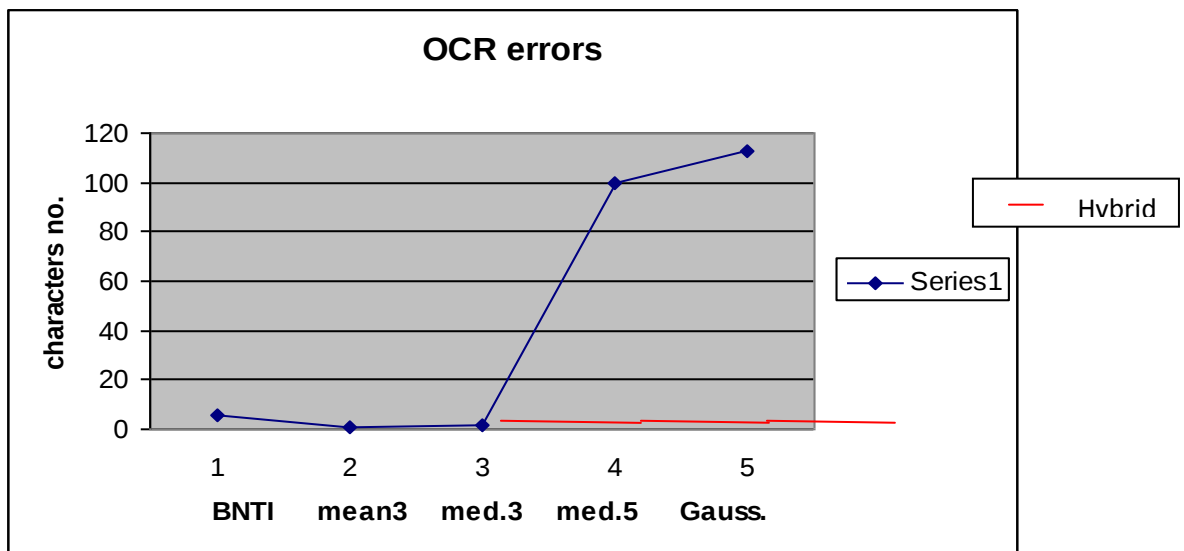


Fig. 13 : Chart 1, "Implementation results of Black noisy text image removal with traditional and Hybrid methods".

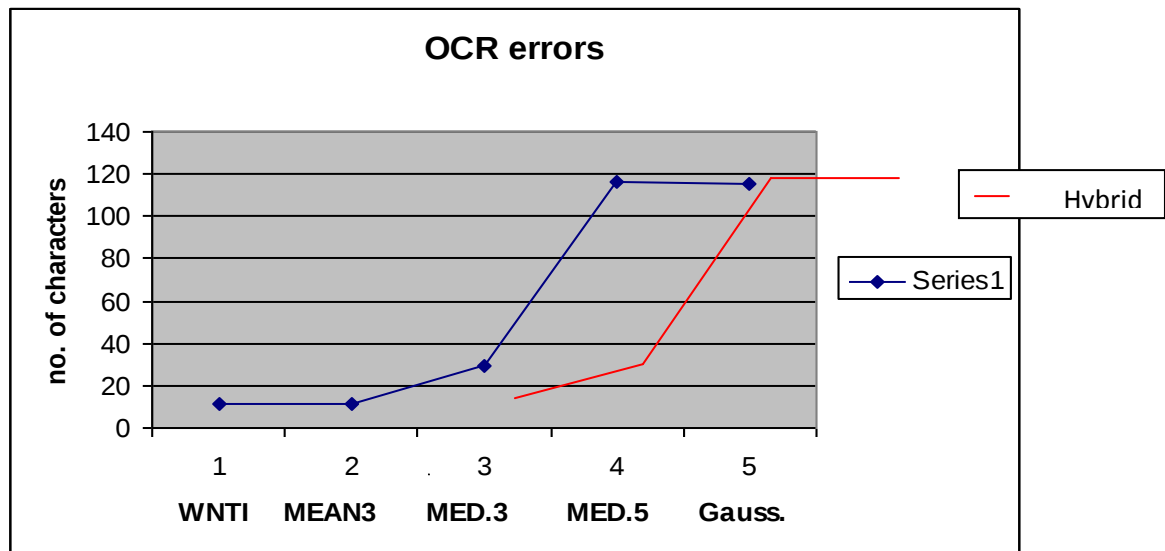


Fig. 13 : Chart 2, "Implementation results of Wight noisy text image removal with traditional and Hybrid methods".

The two charts summarize these results for the sections used in the study. Were In the chart each point shows the number of character unrecognized were the section we used have 126 characters .

CONCLUSION

- 1- applying the a hybrid enhancement filter with BNTI gives results better than with WNTI, the last are poor and many characters show breaks. using a hybrid method with factor=.05 with gaussian filter gives a good result with BNTI removal and with only one error, also with mean filter with factor=.2 gives noise and blurring removal and with only one error.
- 2- The result of the two enhancement methods (median 3*3, mean 3*3) with BNTI according to the no. of characters recognized are the same.
- 3- The result in EWNTI is worse than with EBNTI, and the median3*3 filter gives the best result with BNTI but with the WNTI, the mean3*3 filter gives the best results.
- 4- The above enhancement algorithms work best with image than text image.
- 5- The problem of BNTI is joining between two characters and the problem with WNTI is the gaps that appear in the shape of character.

References

1. Wikipedia, the free encyclopedia,
http://en.wikipedia.org/wiki/Optical_character_recognition.
2. "Reading Machine Speaks Out Loud" , February 1949, Popular Science.
3. Kai Niklas, "Unsupervised Post-Correction of OCR Errors", 2010.

TEST THE EFFECT OF ENHANCEMENT FILTERS ON THE NOISY TEXT IMAGENidhal K. shukur

4. Suen, C.Y., "Future Challenges in Handwriting and Computer Applications", 3rd International Symposium on Handwriting and Computer Applications, Montreal, May 29, 1987.
5. Tappert, Charles C., "The State of the Art in On-line Handwriting Recognition", IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol 12 No 8, August 1990, pp 787-ff.
6. http://finereader.abbyy.com/full_feature_list/ocr_accuracy/
7. Pruslin, Dennis Howard , " Automatic Recognition of sheet Music", 1966.
8. <http://www.visiv.co.uk/quote.htm>
9. Saffady, W. , " Electronic document imaging systems: Design, evaluation, and implementation",1993. Westport, CT: Meckler. 182 p. ISBN 0-88736-840-9.
10. Warner, T. , "Which documents are good candidates for OCR Macworld", 1993, vol.10, no. 11, p. 95.
11. Roger T. Hartley, Kathleen Crumpton, "QUALITY OF OCR FOR DEGRADED TEXT IMAGES".

فحص تأثير خوارزميات التحسين على الصورة النصية المشوشة

ملخص

يتناول هذا البحث سمة مهمة من برامج التعرف الضوئي على الحروف، ويصف انواع برامج المسح الضوئي و المقارنة بينهما ويبين كيف تنفيذ المرشحات على الصورة النصية المشوشة، وذلك باستخدام hybrid method with guassian filter واعطت نتيجة جيدة مع إزالة BNTI، وأيضا عند تنفيذ طريقة hybrid method with mean and median filter تمت ازالة الغواش وإزالة الضبابية . كما أظهرت النتائج أن الطرق التي تم تنفيذها في هذا البحث تعمل بشكل افضل مع EBNTI من نظيرتها WNTI وقللت عدد الاخطاء بشكل واضح .