

GLOBAL JOURNALS

OF COMPUTER SCIENCE AND TECHNOLOGY: F

Graphics & Vision

Impulse Noise Removal
Environment for Robot

Highlight

Web Information Fusion
Analysis of Gray Scale

Dis

VOLUME 13

ISSUE 1

VERSION 1.0



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY: F
GRAPHICS & VISION

GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY: F
GRAPHICS & VISION

VOLUME 13 ISSUE 1 (VER. 1.0)

OPEN ASSOCIATION OF RESEARCH SOCIETY

© Global Journal of Computer Science and Technology. 2013.

All rights reserved.

This is a special issue published in version 1.0 of "Global Journal of Computer Science and Technology" By Global Journals Inc.

All articles are open access articles distributed under "Global Journal of Computer Science and Technology"

Reading License, which permits restricted use. Entire contents are copyright by of "Global Journal of Computer Science and Technology" unless otherwise noted on specific articles.

No part of this publication may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopy, recording, or any information storage and retrieval system, without written permission.

The opinions and statements made in this book are those of the authors concerned. Ultraculture has not verified and neither confirms nor denies any of the foregoing and no warranty or fitness is implied.

Engage with the contents herein at your own risk.

The use of this journal, and the terms and conditions for our providing information, is governed by our Disclaimer, Terms and Conditions and Privacy Policy given on our website <http://globaljournals.us/terms-and-condition/menu-id-1463/>

By referring / using / reading / any type of association / referencing this journal, this signifies and you acknowledge that you have read them and that you accept and will be bound by the terms thereof.

All information, journals, this journal, activities undertaken, materials, services and our website, terms and conditions, privacy policy, and this journal is subject to change anytime without any prior notice.

Incorporation No.: 0423089
License No.: 42125/022010/1186
Registration No.: 430374
Import-Export Code: 1109007027
Employer Identification Number (EIN):
USA Tax ID: 98-0673427

Global Journals Inc.

(A Delaware USA Incorporation with "Good Standing"; Reg. Number: 0423089)

Sponsors: Global Association of Research

Open Scientific Standards

Publisher's Headquarters office

Global Journals Inc., Headquarters Corporate Office,
Cambridge Office Center, II Canal Park, Floor No.
5th, **Cambridge (Massachusetts)**, Pin: MA 02141
United States

USA Toll Free: +001-888-839-7392

USA Toll Free Fax: +001-888-839-7392

Offset Typesetting

Global Association of Research, Marsh Road,
Rainham, Essex, London RM13 8EU
United Kingdom.

Packaging & Continental Dispatching

Global Journals, India

Find a correspondence nodal officer near you

To find nodal officer of your country, please
email us at local@globaljournals.org

eContacts

Press Inquiries: press@globaljournals.org

Investor Inquiries: investers@globaljournals.org

Technical Support: technology@globaljournals.org

Media & Releases: media@globaljournals.org

Pricing (Including by Air Parcel Charges):

For Authors:

22 USD (B/W) & 50 USD (Color)

Yearly Subscription (Personal & Institutional):

200 USD (B/W) & 250 USD (Color)

EDITORIAL BOARD MEMBERS (HON.)

John A. Hamilton, "Drew" Jr.,
Ph.D., Professor, Management
Computer Science and Software
Engineering
Director, Information Assurance
Laboratory
Auburn University

Dr. Henry Hexmoor
IEEE senior member since 2004
Ph.D. Computer Science, University at
Buffalo
Department of Computer Science
Southern Illinois University at Carbondale

Dr. Osman Balci, Professor
Department of Computer Science
Virginia Tech, Virginia University
Ph.D. and M.S. Syracuse University,
Syracuse, New York
M.S. and B.S. Bogazici University,
Istanbul, Turkey

Yogita Bajpai
M.Sc. (Computer Science), FICCT
U.S.A. Email:
yogita@computerresearch.org

Dr. T. David A. Forbes
Associate Professor and Range
Nutritionist
Ph.D. Edinburgh University - Animal
Nutrition
M.S. Aberdeen University - Animal
Nutrition
B.A. University of Dublin- Zoology

Dr. Wenying Feng
Professor, Department of Computing &
Information Systems
Department of Mathematics
Trent University, Peterborough,
ON Canada K9J 7B8

Dr. Thomas Wischgoll
Computer Science and Engineering,
Wright State University, Dayton, Ohio
B.S., M.S., Ph.D.
(University of Kaiserslautern)

Dr. Abdurrahman Arslanyilmaz
Computer Science & Information Systems
Department
Youngstown State University
Ph.D., Texas A&M University
University of Missouri, Columbia
Gazi University, Turkey

Dr. Xiaohong He
Professor of International Business
University of Quinipiac
BS, Jilin Institute of Technology; MA, MS,
PhD,. (University of Texas-Dallas)

Burcin Becerik-Gerber
University of Southern California
Ph.D. in Civil Engineering
DDes from Harvard University
M.S. from University of California, Berkeley
& Istanbul University

Dr. Bart Lambrecht

Director of Research in Accounting and Finance
Professor of Finance
Lancaster University Management School
BA (Antwerp); MPhil, MA, PhD
(Cambridge)

Dr. Carlos García Pont

Associate Professor of Marketing
IESE Business School, University of Navarra
Doctor of Philosophy (Management),
Massachusetts Institute of Technology (MIT)
Master in Business Administration, IESE,
University of Navarra
Degree in Industrial Engineering,
Universitat Politècnica de Catalunya

Dr. Fotini Labropulu

Mathematics - Luther College
University of Regina
Ph.D., M.Sc. in Mathematics
B.A. (Honors) in Mathematics
University of Windsor

Dr. Lynn Lim

Reader in Business and Marketing
Roehampton University, London
BCom, PGDip, MBA (Distinction), PhD,
FHEA

Dr. Mihaly Mezei

ASSOCIATE PROFESSOR
Department of Structural and Chemical
Biology, Mount Sinai School of Medical
Center
Ph.D., Eötvös Loránd University
Postdoctoral Training,
New York University

Dr. Söhnke M. Bartram

Department of Accounting and Finance
Lancaster University Management School
Ph.D. (WHU Koblenz)
MBA/BBA (University of Saarbrücken)

Dr. Miguel Angel Ariño

Professor of Decision Sciences
IESE Business School
Barcelona, Spain (Universidad de Navarra)
CEIBS (China Europe International Business School).
Beijing, Shanghai and Shenzhen
Ph.D. in Mathematics
University of Barcelona
BA in Mathematics (Licenciatura)
University of Barcelona

Philip G. Moscoso

Technology and Operations Management
IESE Business School, University of Navarra
Ph.D in Industrial Engineering and
Management, ETH Zurich
M.Sc. in Chemical Engineering, ETH Zurich

Dr. Sanjay Dixit, M.D.

Director, EP Laboratories, Philadelphia VA
Medical Center
Cardiovascular Medicine - Cardiac
Arrhythmia
Univ of Penn School of Medicine

Dr. Han-Xiang Deng

MD., Ph.D
Associate Professor and Research
Department Division of Neuromuscular
Medicine
Davee Department of Neurology and Clinical
Neuroscience
Northwestern University
Feinberg School of Medicine

Dr. Pina C. Sanelli

Associate Professor of Public Health
Weill Cornell Medical College
Associate Attending Radiologist
NewYork-Presbyterian Hospital
MRI, MRA, CT, and CTA
Neuroradiology and Diagnostic
Radiology
M.D., State University of New York at
Buffalo, School of Medicine and
Biomedical Sciences

Dr. Roberto Sanchez

Associate Professor
Department of Structural and Chemical
Biology
Mount Sinai School of Medicine
Ph.D., The Rockefeller University

Dr. Wen-Yih Sun

Professor of Earth and Atmospheric
SciencesPurdue University Director
National Center for Typhoon and
Flooding Research, Taiwan
University Chair Professor
Department of Atmospheric Sciences,
National Central University, Chung-Li,
TaiwanUniversity Chair Professor
Institute of Environmental Engineering,
National Chiao Tung University, Hsin-
chu, Taiwan.Ph.D., MS The University of
Chicago, Geophysical Sciences
BS National Taiwan University,
Atmospheric Sciences
Associate Professor of Radiology

Dr. Michael R. Rudnick

M.D., FACP
Associate Professor of Medicine
Chief, Renal Electrolyte and
Hypertension Division (PMC)
Penn Medicine, University of
Pennsylvania
Presbyterian Medical Center,
Philadelphia
Nephrology and Internal Medicine
Certified by the American Board of
Internal Medicine

Dr. Bassey Benjamin Esu

B.Sc. Marketing; MBA Marketing; Ph.D
Marketing
Lecturer, Department of Marketing,
University of Calabar
Tourism Consultant, Cross River State
Tourism Development Department
Co-ordinator , Sustainable Tourism
Initiative, Calabar, Nigeria

Dr. Aziz M. Barbar, Ph.D.

IEEE Senior Member
Chairperson, Department of Computer
Science
AUST - American University of Science &
Technology
Alfred Naccash Avenue – Ashrafieh

PRESIDENT EDITOR (HON.)

Dr. George Perry, (Neuroscientist)

Dean and Professor, College of Sciences

Denham Harman Research Award (American Aging Association)

ISI Highly Cited Researcher, Iberoamerican Molecular Biology Organization

AAAS Fellow, Correspondent Member of Spanish Royal Academy of Sciences

University of Texas at San Antonio

Postdoctoral Fellow (Department of Cell Biology)

Baylor College of Medicine

Houston, Texas, United States

CHIEF AUTHOR (HON.)

Dr. R.K. Dixit

M.Sc., Ph.D., FICCT

Chief Author, India

Email: authorind@computerresearch.org

DEAN & EDITOR-IN-CHIEF (HON.)

Vivek Dubey(HON.)

MS (Industrial Engineering),

MS (Mechanical Engineering)

University of Wisconsin, FICCT

Editor-in-Chief, USA

editorusa@computerresearch.org

Sangita Dixit

M.Sc., FICCT

Dean & Chancellor (Asia Pacific)

deanind@computerresearch.org

Suyash Dixit

(B.E., Computer Science Engineering), FICCTT

President, Web Administration and

Development , CEO at IOSRD

COO at GAOR & OSS

Er. Suyog Dixit

(M. Tech), BE (HONS. in CSE), FICCT

SAP Certified Consultant

CEO at IOSRD, GAOR & OSS

Technical Dean, Global Journals Inc. (US)

Website: www.suyogdixit.com

Email: suyog@suyogdixit.com

Pritesh Rajvaidya

(MS) Computer Science Department

California State University

BE (Computer Science), FICCT

Technical Dean, USA

Email: pritesh@computerresearch.org

Luis Galárraga

J!Research Project Leader

Saarbrücken, Germany

CONTENTS OF THE VOLUME

- i. Copyright Notice
 - ii. Editorial Board Members
 - iii. Chief Author and Dean
 - iv. Table of Contents
 - v. From the Chief Editor's Desk
 - vi. Research and Review Papers
-
- 1. Impulse Noise Removal from Digital Images- A Computational Hybrid Approach. *1-7*
 - 2. Analysis of Cross-Media Web Information Fusion for Text and Image Association- A Survey Paper. *9-15*
 - 3. Edge Identification During Fire Environment for Robot. *17-19*
 - 4. Performance Analysis of Gray Scale and Color Iris with Multidomain Feature Normalization and Dimensionality Reduction. *21-29*
-
- vii. Auxiliary Memberships
 - viii. Process of Submission of Research Paper
 - ix. Preferred Author Guidelines
 - x. Index



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
GRAPHICS & VISION

Volume 13 Issue 1 Version 1.0 Year 2013

Type: Double Blind Peer Reviewed International Research Journal

Publisher: Global Journals Inc. (USA)

Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Impulse Noise Removal from Digital Images- A Computational Hybrid Approach

By M.Jayamanmadharao, M S Ramanaidu & Dr. KVVS Reddy

Andhra University, A.P., India

Abstract - In digital Image Processing, removal of noise is a highly demanded area of research. Impulsive noise is common in images which arise at the time of image acquisition and or transmission of images. In this paper, a new hybrid filtering algorithm is presented for the removal of impulse noise from digital images. Here, we replace the impulse noise corrupted pixel by the median of the pixel scanned in four directions. The experimental results of this filter applied on various images corrupted with almost all ratios of impulse noise favor the filter in terms of objectivity than many of the other prominent impulse noise filters.

GJCST-F Classification: I.4.1



Strictly as per the compliance and regulations of:



Impulse Noise Removal from Digital Images- A Computational Hybrid Approach

M.Jayamanmadharao^α, M S Ramanaidu^σ & Dr. KWVS Reddy^ρ

Abstract - In digital Image Processing, removal of noise is a highly demanded area of research. Impulsive noise is common in images which arise at the time of image acquisition and or transmission of images. In this paper, a new hybrid filtering algorithm is presented for the removal of impulse noise from digital images. Here, we replace the impulse noise corrupted pixel by the median of the pixel scanned in four directions. The experimental results of this filter applied on various images corrupted with almost all ratios of impulse noise favor the filter in terms of objectivity than many of the other prominent impulse noise filters.

I. INTRODUCTION

Digital images play an important role both in daily life applications such as Satellite television, Magnetic Resonance Imaging, Computer Tomography as well as in areas of research and technology such as geographical information systems and astronomy. Digital images are often corrupted by different types of noise during its acquisition and transmission phase. Such degradation negatively influences the performance of many image processing techniques and a preprocessing module to filter the images is often required [1, 2]. To enhance the quality of images various images enhancement or restoration techniques are use. Efficiency of every method is depending on the quality of input images.

The overall noise characteristics in an image depend on many factors, including the type of sensor, pixel dimensions, temperature, exposure time, and speed. The goal of image denoising is to remove the noise while retaining the important signal features. The denoising of a natural image corrupted by noise is an important problem in image processing. Image denoising still remains a challenge for researchers because noise removal introduces artifacts and causes blurring of the images. This paper discusses the methods of noise reduction of impulse noise image using hybrid approach [3-5].

II. LITERATURE REVIEW

The one of the emerging field of image processing is removal of noise from a contaminated image. Many researchers have suggested a large

number of algorithms and compared their results. The main thrust on all such algorithms is to remove impulsive noise while preserving image details. Some schemes utilize detection of impulsive noise followed by filtering where as others filter all the pixels irrespective of corruption. In this section an attempt has been made for a literature review for the filtering of random-valued impulsive noise.

Umesh Ghanekar et.al [6] presents a new filtering scheme based on contrast enhancement within the filtering window for removing the random valued impulse noise. The application of a nonlinear function for increasing the difference between a noise-free and noisy pixels results in efficient detection of noisy pixels. As the performance of a filtering system, in general, depends on the number of iterations used, an effective stopping criterion based on noisy image characteristics to determine the number of iterations is also proposed. Extensive simulation results exhibit that the proposed method significantly outperforms many other well-known techniques. The performance of the proposed scheme has been compared with many existing techniques. The efficacy of the proposed method is demonstrated by extensive simulations. The experimental results exhibit significant improvement in the performance over several other methods. Also, the proposed method requires less iteration in comparison with some recently proposed methods.

Yiqiu Dong et.al [7] proposes an image statistic for detecting random-valued impulse noise. By this statistic, it can identify most of the noisy pixels in the corrupted images. Combining it with an edge-preserving regularization, it obtain a powerful two-stage method for denoising random-valued impulse noise, even for noise levels as high as 60%. Simulation results show that proposed method is significantly better than a number of existing techniques in terms of image restoration and noise detection. The authors say that for the other methods, there are still some noticeable noise unremoved and there exist some loss and discontinuity of the details, such as the hair around the mouth of the baboon and the edges of the bridge. In contrast, the visual qualities of restored images are quite good, even with the abundance of image details and the high noise level present in the images. Simulation results show that proposed method outperforms a number of existing methods both visually and quantitatively.

Author α : Professor, Department of EIE, AITAM, Tekkali, A.P., India.

Author σ : Associate Professor, Department of EIE, AITAM, Tekkali, A.P., India.

Author ρ : Professor, Department of ECE, AU College of Engineering, Andhra University, A.P., India.

Stefan Schulte et.al [8] presented a new impulse noise reduction method for color images. Color images that are corrupted with impulse noise are generally filtered by applying a grayscale algorithm on each color component separately or using a vector-based approach where each pixel is considered as a single vector. It discusses an alternative technique which gives a good noise reduction performance while much less artefacts are introduced. The main difference between the proposed method and other classical noise reduction methods is that the color information is taken into account to develop a better impulse noise detection method and a noise reduction method that filters only the corrupted pixels while preserving the color and the edge sharpness.

Abdul Majid et.al [9] proposes a novel impulse noise removal scheme that emphasizes on few noise-free pixels and small neighborhood. The proposed scheme searches noise-free pixels within a small neighborhood. This scheme has provided better performance as compared to existing approaches. Moreover, this scheme is capable to restore the corrupted images while preserving edges and fine details. Experimental results show that the proposed scheme is capable of removing impulse noise effectively while preserving the fine image details. Especially, this approach has shown effectiveness against high impulse noise density.

Leah Bar et.al [10] presents a unified variational approach to image deblurring and impulse noise removal. The objective functional consists of a fidelity term and a regularizer. Data fidelity is quantified using the robust modified L1 norm, and elements from the Mumford-Shah functional are used for regularization. It shows that the Mumford-Shah regularizer can be viewed as an extended line process. It reflects spatial organization properties of the image edges that do not appear in the common line process or anisotropic diffusion. This allows distinguishing outliers from edges and leads to superior experimental results.

Zayed M. Ramadan [11] proposes a two-stage adaptive method for restoration of images corrupted with impulse noise. In the first stage, the pixels which are most likely contaminated by noise are detected based on their intensity values. In the second stage, an efficient average filtering algorithm is used to remove those noisy pixels from the image. Only pixels which are determined to be noisy in the first stage are processed in the second stage. The remaining pixels of the first stage are not processed further and are just copied to their corresponding locations in the restored image. The experimental results for the proposed method demonstrate that it is faster and simpler than even median filtering, and it is very efficient for images corrupted with a wide range of impulse noise densities varying from 10% to 90%. Because of its simplicity, high speed, and low computational complexity, the proposed

method can be used in realtime digital image applications, e.g., in consumer electronic products such as digital televisions and cameras.

Jian-Feng Cai et.al [12] proposes a two-phase approach to restore images corrupted by blur and impulse noise. In the first phase, It identifies the outlier candidates—the pixels that are likely to be corrupted by impulse noise. It considers that the remaining data pixels are essentially free of outliers. Then in the second phase, the image is deblurred and denoised simultaneously by a variational method by using the essentially outlier-free data. In general the two-phase method for salt-and-pepper noise performs better than for random-valued noise: it can handle salt-and-pepper noise as high as 90% but random-valued noise for about 55%. The main reason is that the former is easier to detect than the latter in the first phase. In fact, AMF is a very good detector for salt-and-pepper noise, and almost all the noise positions can be detected even when the noise ratio is very high. In addition, with salt-and-pepper noise, most of the noisy pixels are much more dissimilar to regular pixels, hence are easier to detect. However, there is no good detector for random-valued noise when the noise ratio is high. The performance for random-valued noise can be improved if a better noise detector can be found in the first phase.

R K Kulkarni et.al [13] proposes a simple yet effective algorithm for effectively denoising the extremely corrupted image by impulse noise. The proposed method first classifies the pixels into two classes, which are “noise free pixel” and “noisy pixel” based on the intensity values. The corrupted pixels are replaced by “alpha trim- mean” value of uncorrupted pixels in the filtering window. The method adaptively changes the size of filtering window based on the number of “noise free pixels”. Because of this, the proposed method removes the noise much more effectively even at noise density as high as 90% and preserves the good image quality. Experimental results show that this method always produces good output, even when tested for high level of noise. The details inside the image are preserved. This method is simple and relatively a fast method and suitable for consumer electronic products such as digital camera.

X D Jiang [14] proposed a new nonlinear filter for attenuating impulsive noise while preserving image details. The filter truncates the grey value of a pixel to the maximal or minimal value of its enclosed surrounding band. Impulsive noise inside the band is thus attenuated while image details are preserved as long as they stretch to the band. The recursive form of the proposed filter leads to a simple architecture for fast implementation. Theoretical analysis and experimental results demonstrate the effectiveness of this new filter for both noise attenuation and detail preservation. According to simulation results the proposed filter outperforms the standards median filter, the centre-

weighted median filter and the unidirectional multistage median filter in terms of mean absolute error and filtering speed.

J.Harikiran et.al [15] introduces the concept of image fusion of filtered noisy images for impulse noise reduction. Image fusion is the process of combining two or more images into a single image while retaining the important features of each image. Multiple image fusion is an important technique used in military, remote sensing and medical applications. Five different filtering algorithms are used individually for filtering the image captured from the sensor. The filtered images are fused to obtain a high quality image compared to individually denoised images. The performance of filters was evaluated by computing the mean square error (MSE) between the original image and filtered image. Experimental results show that this method is capable of

producing better results compared to individually denoised images.

III. PROPOSED METHOD

The main challenge in impulse noise removal is to suppress the noise as well as to preserve the details (edges). This paper presents a simple & effective way to remove the impulse and random noises from the digital image. The first step is to detect the impulse and random noise from the image as shown in Fig 1. In this stage, based on the only intensity value, the pixels are roughly divided into two classes, which are “noise-free pixel” and “noise-pixel”. Then the second stage is to eliminate the impulse and the random noise from the image.

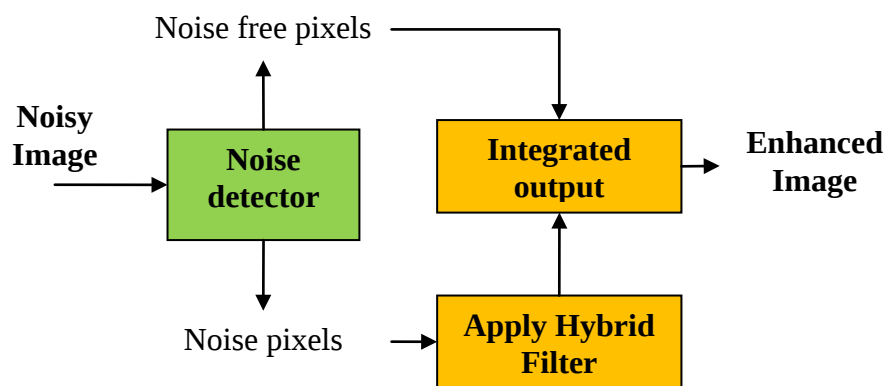


Fig. 1 : System Structure

The proposed filter operates on impulse noise densities without jeopardizing image fine details and textures. Fast and automated algorithm is focused. The proposed filter does not require any tedious tuning or time consuming training of parameters as well. No priori threshold is to be given. Instead, the threshold is computed locally from image pixels intensity values in a sliding window using weighted statistics. More precisely, the weighted mean value and the weighted standard deviation are estimated in the current window. The weights are the inverse of the distance between the weighted mean value of pixels in a given window and the considered pixel. A result is that impulse noise does not corrupt the determination of these statistics from which the Threshold is derived. Noise-free pixels are relatively easy to be selected by utilizing the binary decision. A limit for window is set to contain a minimum number of pixels avoid loss of image details. In filtering mechanism, the proposed filter adopts fuzzy reasoning to deal with uncertainties present in the local information. These uncertainties, e.g. thin lines or pixels at edges being mistaken as noise-pixels, are caused by the nonlinear nature of impulse noise. The fuzzy set is processed by calculating local information to produce a suitable fuzzy membership value.

a) Proposed Algorithm

In an image contaminated by random-valued impulse noise, the detection of noisy pixel is more difficult in comparison with fixed valued impulse noise, as the gray value of noisy pixel may not be substantially larger or smaller than those of its neighbors. Due to this reason, the conventional median-based impulse detection methods do not perform well in case of random valued impulse noise. The numerical Threshold value is defined a priori or chosen after many data dependant tests. The literature shows that an optimal threshold in the sense of the mean square error can be obtained for most real data. However, Threshold suitable for a particular image is not necessarily adapted to another one. To overcome this problem, the following algorithm is proposed.

Step 1:	Read the input image and add Random Valued Impulse Noise to the image.
Step 2:	Compute the weighted mean value of the window. $M(i, j) = \frac{\sum_{m,n} w_{m,n} x_{i+m,j+n}}{\sum_{m,n} w_{m,n}} \quad (1)$
Step 3:	Weighted standard deviation is calculated using the weighted mean value.

$$\sigma(i, j) = \frac{\sum_{m,n} w_{m,n} (X_{i+m,j+n} - M(i,j))^2}{\sum_{m,n} w_{m,n}} \quad (2)$$

Step 4: Threshold is obtained from the above statistical parameters which is given by $\alpha\sigma(i,j)$, $\alpha=1$

Step 5: Noisy pixel is found when difference between centre pixel and weighted mean exceeds threshold.

Step 6: Binary flag represents as follows:

1- Noisy pixel
0- Noise free pixel.

Step 7: Compute the median value for noise free pixels

Step 8: Determine absolute difference.

$D(i,j) = \max\{d(m,n)\} \text{-----}(3)$

Where $d(m,n) = \frac{|x(m,n) - x(i,j)|}{255}$ $x(n,m) \in \omega(x)_{(i,j)}$

$X(i,j)$ is the centre pixel in window.

Step 9: Compute the fuzzy membership value $F(i,j)$

$F(i,j) = 0$ if $D(i,j) < T_1$

$F(i,j) = 1$ if $D(i,j) > \text{or equal to } T_2$

Step 10: Compute the restoration term $y(i,j)$

IV. PERFORMANCE EVALUATION

There are two main methodologies are used to estimate the quality of images. They are subjective evaluation and objective evaluation.

a) Subjective Evaluation

To obtain reliable quality rating, subjective viewing tests carried on post processed images. The rating for given image is given as excellent, good, average, bad etc. but the result of given rating depends on the following factors.

- The experience and motivation of the subject.
- The range of the picture used,
- The conditions under which the pictures are viewed

b) Objective quality measures

The simplest and still most commonly used objective measures are Mean Square Error (MSE) and the Peak Signal to Noise Ratio (PSNR). The mathematical formulae for the two are

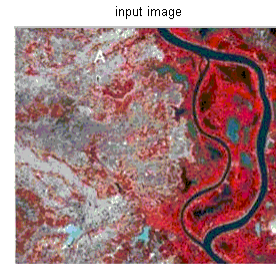
$$MSE = \frac{1}{MN} \sum_{y=1}^M \sum_{x=1}^N [I(x,y) - I^1(x,y)]^2$$

$$PSNR = 20 * \log_{10} (255 / \sqrt{MSE})$$

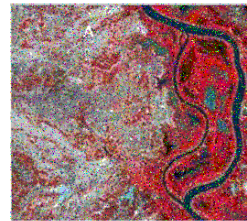
Where $I(x,y)$ is the original image, $I^1(x,y)$ is the approximated version and M,N are the dimensions of the images. These measures give simple mathematical deviation between original image and reconstructed image. They operate solely on pixel by pixel basis.

V. RESULTS

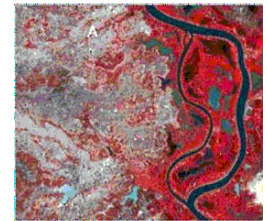
The designed filter was tested with impulse noise in remote sensing images under several noise conditions, it shown in Fig 2.



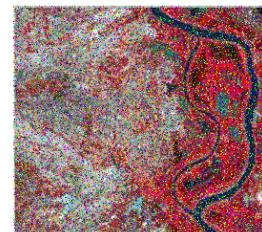
Noise image(10%)



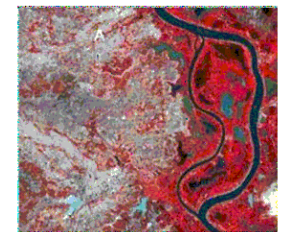
Filtered image



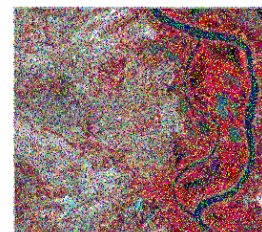
Noise image(20%)



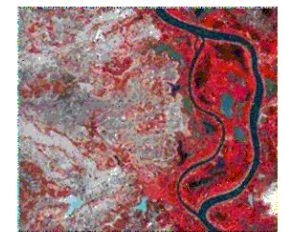
Filtered image



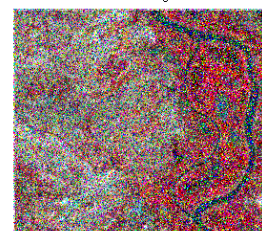
Noise image(30%)



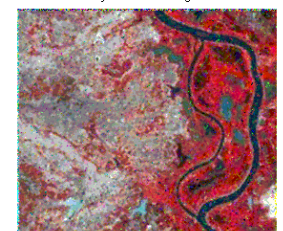
Filtered image



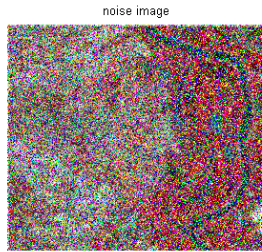
Noise image(40%)



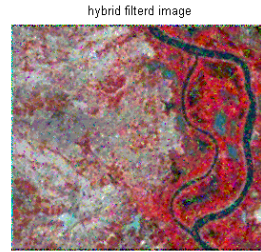
Filtered image



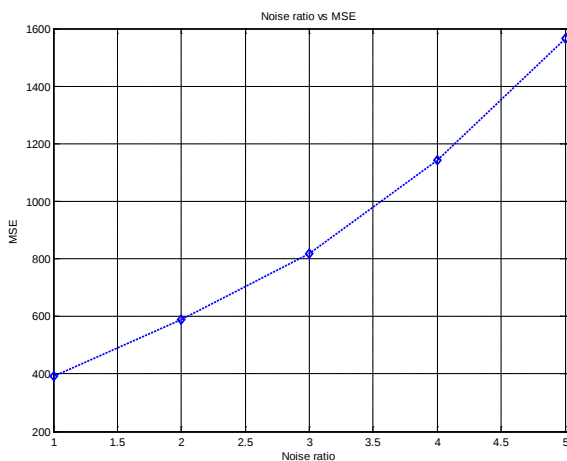
Noise image(50%)



Filtered image

*Figure 2* : Original, noise and filtered images of hybrid method*Table 1* : Noise ratio, MSE and PSNR

Noise (%)	MSE	PSNR
10	392.7394	44.379
20	589.0978	40.857
30	818.9550	37.996
40	1.1439e+00	35.094
50	1.5667e+00	32.362

*Figure 3* : Shows the Noise ratio vs MSE

These results are shown in Table 1. The original image is corrupted by mixed noise. Based on the result data we construct the graphs for MSE, PSNR under different noise ratios, which are shown in Fig 3, 4. From these figures we said that Image with lower MSE and a high PSNR, it means the image is a better one.

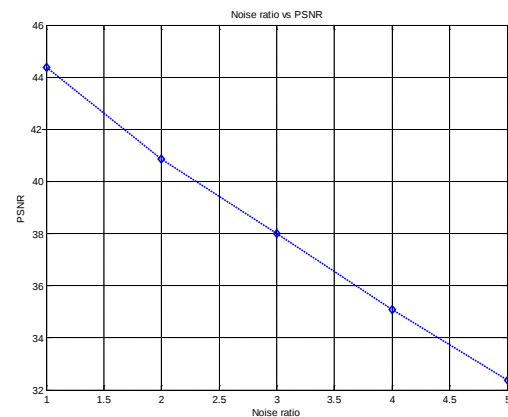
*Figure 4* : Noise ratio vs PSNR

Table 2 shows the MSE, PSNR values of the proposed method which is compared with the other methods. Based on this data we construct the graph as shown in Fig 5, from this graph we clearly observe that the proposed method performs better in the form of high MSE with low PSNR.

Table 2 : Performance evaluation with different filtered images at same noise level for Remote sensing image

Filter	Noise level 10%		Noise level 20%		Noise level 30%		Noise level 40%		Noise level 50%	
	MSE	PSNR	MSE	PSNR	MSE	PSNR	MSE	PSNR	MSE	PSNR
Vector median filter(VMF)	47.984	31.23	86.9	28.6	146	26.6	333.88	24	696	19.6
Rank conditioned VMF	48.368	31.4	83.4	29.2	142	26.9	312	23.4	683	19.7
Rank conditioned & threshold VMF	68.7	30.6	99.6	30.12	184	23	411	20.3	918	18.6
Center weighted VMF	173.42	26.1	273	22.9	482	21	944	18	2014	14.9
Absolute deviation VMF	19.58	34.2	28.27	31.8	77.9	29	211	22	346.0	22.8
Spatial Median Filter Image	12.88	37.032	16.86	35.86	30.64	33.26	75.816	29.369	174	25.724
Modified Spatial Median Filter	13.99	38.86	14.76	36.78	27.44	36.321	72.89	31.789	162	22.346
Proposed filter(Hybrid)	392.7394	44.3795	589.0978	40.8579	818.9550	37.9964	1.1439e+003	35.0940	1.5667e+003	32.3621

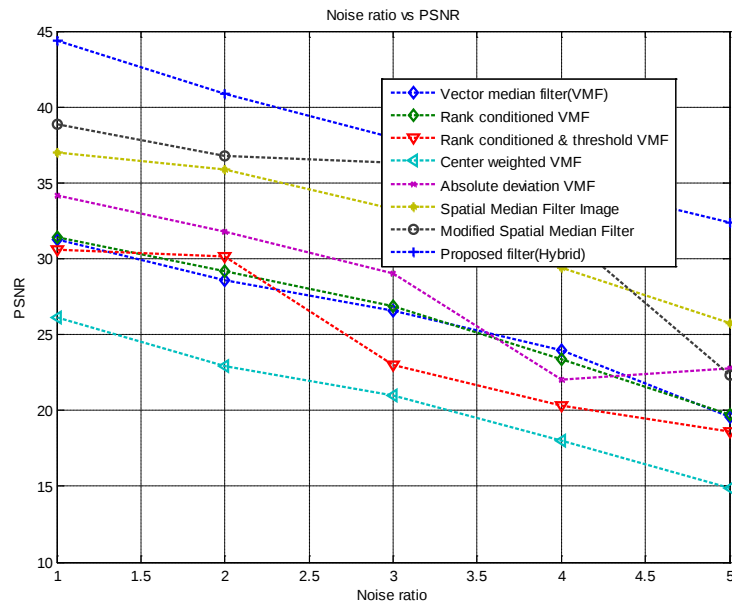


Figure 5 : Shows the comparison among the filters

VI. CONCLUSIONS

A novel hybrid filtering operator for removing mixed noise from digital images is presented. The fundamental superiority of the proposed operator over most other operators is that it efficiently removes impulse noise from digital images while preserving thin lines and edges in the original image. Extensive simulation results verify its excellent impulse detection and detail preservation abilities by attaining the highest PSNR and lowest MAE values across a wide range of noise densities. Thus rampant loss of image is reduced without jeopardizing image fine details.

REFERENCES RÉFÉRENCES REFERENCIAS

1. M. Tulin Yıldırım, Alper Bas_turk and M. Emin Yuksel, Impulse Noise Removal From Digital Images by a Detail-Preserving Filter Based on Type-2 Fuzzy Logic, IEEE transactions on fuzzy systems, pp-920-928, 2008.
2. Tom Mélangé, Mike Nachtegaal, and Etienne E. Kerre, Fuzzy Random Impulse Noise Removal from Color Image Sequences, IEEE transactions on image processing, pp-959-970, 2011.
3. Wenbin Luo, Efficient Removal of Impulse Noise from Digital Images, IEEE Transactions, pp.523-527, 2006.
4. Zhou Wang and David Zhang, Progressive Switching Median Filter for the Removal of Impulse Noise from Highly Corrupted Images, IEEE transactions on circuits and systems—II: analog and digital signal processing, pp.78-80,1999.
5. Roman Garnett, Timothy Huegerich, Charles Chui, and Wenjie He, A Universal Noise Removal Algorithm with an Impulse Detector, IEEE transactions on image processing, pp. 1747-1754, 2005.
6. Umesh Ghanekar, Awadhesh Kumar Singh, and Rajoo Pandey, A Contrast Enhancement-Based Filter for Removal of Random Valued Impulse Noise, IEEE Transactions on signal processing letters, pp-47-50, 2010.
7. Yiqiu Dong, Raymond H. Chan, and Shufang Xu, A Detection Statistic for Random-Valued Impulse Noise, IEEE transactions on image processing, pp. 1112-1120, 2007.
8. Stefan Schulte, Samuel Morillas, Valentín Gregori, and Etienne E. Kerre, A New Fuzzy Color Correlated Impulse Noise Reduction Method, IEEE transactions on image processing, pp.2565-2575, 2007.
9. Abdul Majid ,Choong-Hwan Lee ,Muhammad ,Tariq Mahmood ,Tae-Sun Choi, Impulse noise filtering based on noise-free pixels using genetic programming, Knowl Inf Syst, Springer-Verlag, pp.505-526, 2012.
10. Leah Bar, Nahum Kiryati, Nir Sochen, Image Deblurring in the Presence of Impulsive Noise, International Journal of Computer Vision, Springer Science and Business Media, pp. 279-298, 2006.
11. Zayed M. Ramadan, Efficient Restoration Method for Images Corrupted with Impulse Noise, Circuits Syst Signal Process , Springer Science and Business Media, LLC, pp.1397-1406, 2012.
12. Jian-Feng Cai •, Raymond H. Chan •, Mila Nikolova, Fast Two-Phase Image Deblurring Under Impulse Noise, J Math Imaging Vis , Springer Science and Business Media, LLC, pp. 46-53, 2010
13. R.K.Kulkarni, S.Meher, J.M.Nair, An Adaptive Switching Filter for Removing Impulse Noise from

Highly corrupted Images, ICGST-GVIP Journal, pp.47-53, October 2010.

14. X D Jiang, Image detail-preserving filter for impulsive noise attenuation, IEEE proceedings on Visual image signal process, pp.179-185, 2003.
15. J.Harikiran B.Saichandana B.Divakar, Impulse Noise Removal in Digital Images, International Journal of Computer Applications (IJCA), p. 39-4., 2010.



This page is intentionally left blank



Analysis of Cross-Media Web Information Fusion for Text and Image Association- A Survey Paper

By M. Priyanka, B.Sunita Devi, S.M. Riyazoddin & M. Janga Reddy

CMR Institute of Technology, Hyderabad, Andhra Pradesh

Abstract - The image comprises of the text- and content-based features. Images can be represented using both text-and content-based features. Web information fusion can be defined as the problem of collating and tracking information related to specific topics on the World Wide Web. But the main concern is image and text association, a cornerstone of cross-media web information fusion. Two methods have been described .The first method based on vague transformation measures the information similarity between the visual features and the textual features through a set of predefined domain-specific information categories. Another method uses a neural network to learn direct mapping between the visual and textual features by automatically and incrementally summarizing the associated features into a set of information templates. Despite their distinct approaches, our experimental results on a terrorist domain document set show that both methods are capable of learning associations between images and texts from a small training data set. Another method encompasses a variety of techniques relating to document summarization and text- and content-based image retrieval.

Keywords : *text, image, image processing, image and text association, fusion, context based image representation, search engine.*

GJCST-F Classification: *1.5.0*



Strictly as per the compliance and regulations of:



Analysis of Cross-Media Web Information Fusion for Text and Image Association - A Survey Paper

M. Priyanka ^a, B.Sunita Devi ^a, S.M. Riyazoddin ^a & M. Janga Reddy ^a

Abstract - The image comprises of the text- and content-based features. Images can be represented using both text-and content-based features. Web information fusion can be defined as the problem of collating and tracking information related to specific topics on the World Wide Web. But the main concern is image and text association, a cornerstone of cross-media web information fusion. Two methods have been described .The first method based on vague transformation measures the information similarity between the visual features and the textual features through a set of predefined domain-specific information categories. Another method uses a neural network to learn direct mapping between the visual and textual features by automatically and incrementally summarizing the associated features into a set of information templates. Despite their distinct approaches, our experimental results on a terrorist domain document set show that both methods are capable of learning associations between images and texts from a small training data set. Another method encompasses a variety of techniques relating to document summarization and text- and content-based image retrieval. The text-based approaches described utilize the Unified Medical Language System (UMLS) synonymy to identify concepts in information requests and image-related text in order to re-trieve semantically relevant images. Our image content-based approaches utilize similarity metrics based on computed "visual concepts" and low- level image features to identify visually similar images. The edge model for image and text association detects sharp changes in image brightness and captures important events and changes in properties of the world. This paper discuss about the analysis of the previous work done so that we can suggest an enhanced fusion method for text and Image association.

Keywords : text, image, image processing, image and text association, fusion, context based image representation, search engine.

1. INTRODUCTION

The diverse and distributed nature of the information published on the World Wide Web has made it difficult to collate and track information related to specific topics. Although web search engines have reduced information overloading to a certain extent, the information in the retrieved documents still contains a lot of redundancy. Techniques are needed in web information fusion, involving filtering of irrelevant and redundant information, collating of information

according to themes, and generation of coherent presentation. Document summarization is a commonly based technique used for information fusion and has been a topic of discussion among the scientists. However, most document summarization techniques focus only on text [11], [12], [13].documents whereas in today's world the documents on the net now consist of non text content such as images, video, and sound etc. thus collating multimedia information poses a great challenge in the web information fusion. The image comprises of the text- and content-based features. Images can be represented using both text-and content-based features. The purpose of detecting sharp changes in image brightness is to capture important events and changes in properties of the world. An image-text association refers to a pair of image and text segment that is semantically related to each other in a web page. A sample image-text association is shown in Fig. 1. Identifying such associations enables one to provide a coherent multimedia summarization of the web documents.



Figure 1 : An associated image text pair

The image-text learning association is similar to the task of automatic annotation but has important differences [15]. Whereas image annotation concerns annotating images using a set of predefined keywords, image-text association links images to free text segments in natural language. The methods for image annotation are thus not directly applicable to the problem of identifying image-text associations.

A key issue of using a machine learning approach to image-text associations is the lack of large

^a Author : Department of computer Science and Engineering CMR Institute of Technology, Hyderabad, Andhra Pradesh.

training data sets. Referring to Fig. 2, two associated image-text pairs (I-T pairs) not only share partial visual (smokes and fires) and textual features ("attack") but also have different visual and textual contents. As the two I-T pairs are actually on similar topics (describing scenes of terror attacks), the distinct parts, such as the visual content ("black smoke" which can be represented using low-level color and texture features) of the image in I-T pair 2 and the term "Blazing" (underlined> in I-T pair 1, could be potentially associated. Such useful associations, which convey the information patterns in the domain but are not represented by the training data set, implicit associations.

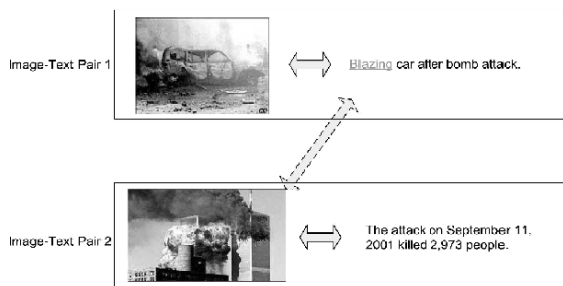


Figure 2 : An image-text Pair

II. LEARNING ASSOCIATION RULES FOR IMAGE CONTENTS AND SEMANTIC CONCEPTS

Association rule mining (ARM) [14] is originally used for discovering association patterns in transaction databases. An association rule is an implication of the form $X \Rightarrow Y$, where X, Y is a subset of I (called item sets or patterns) and $X \cap Y$ is empty. In the domain of market-basket analysis, such an association rule indicates that the customers who buy the set of items X are also likely to buy the set of items Y . Mining association rule from multimedia data is usually a straightforward extension of ARM in transaction databases.

A pixel-based approach presented by Ding et al. [16] is based to deduce associations between pixels' spectral features and semantic concepts. In the model each pixel is treated as a transaction, while the set ranges of the pixel's spectral bands and auxiliary concept labels (e.g., crop yields) are considered as items. However, it was pointed out that using individual pixel as transaction may lose the context information of surrounding locations which are usually very useful for determine the image semantics. This motivated them to use image and rectangular image regions as transactions and items.

III. IMAGE ANNOTATION USING STATISTICAL MODEL

A major deficiency of the existing machine learning and statistic-based automatic multimedia annotation methods is that they usually assign a fixed

number of keywords or concepts to a media object. Therefore, there will inevitably be some media objects assigned with unnecessary annotations and some others assigned with insufficient annotations. In addition, image annotation typically uses a relatively small set of domain-specific key terms (class labels, concepts, or categories) for labeling the images. Our task of discovering semantic image-text associations from Web documents however do not assume such a set of selected domain specific key terms. In fact, although our research focuses on image and text contents from the terrorist domain, both domain-specific and general-domain information (e.g., general-domain terms "people" or "bus") are incorporated in our learning paradigm. With the above considerations, it is clear that the existing image annotation methods are not directly applicable to the task of image-text association. Furthermore, model evolution is not well supported by the above methods. Specifically, after the machine learning and statistic models are trained, they are difficult to update. Moreover, the above methods usually treat the semantic concepts in media objects as separate individuals without considering relationships between the concepts for multimedia content representation and annotation.

IV. IMAGE REPRESENTATION

The image comprises of the text- and content-based features. Images can be represented using both text-and content-based features. Text-based features include text that pertains to an image, such as in captions and "mentions" (snippets of text within the body of an article that discuss an image), and content-based features include information derived from the image itself, such as shapes, colors and textures. The text- and content-based image representation is described below.

a) Text-Based Features

Each image in the ImageCLEFmed'10 collection [7] is represented as a structured document of image-related text. The representation includes the title, abstract, and MeSH terms¹ of the article in which the image appears as well as the image's caption and mention. The content of an image's caption is organized into the well-formed clinical question framework following the method described by Demner - Fushman and Lin [2]. Extractors identify UMLS concepts related to problems, interventions, age, anatomy, drugs, and image modality. Textual Regions of Interest (ROIs) is extracted from image captions. A textual ROI is a noun phrase describing the content of an interesting region of an image which is identified within a caption by a pointer. For example, in the caption "MR image reveals hypo intense indeterminate nodule (arrow)," the word arrow points to the ROI containing a hypo intense indeterminate nodule. The above structured documents can be indexed and searched with a traditional search

engine or the extracted concepts may be combined with additional features (discussed below) for use in a multimodal representation. "Keywords" in a structured document D_j can be represented as an N-dimensional feature vector also.

$$f_j^{\text{keyword}} = [w_{j1}, w_{j2}, \dots, w_{jN}]^T$$

Where w_{jk} denotes the weight (typically tf-idf) of keyword tk in document D_j .

b) Image Content-Based Features

In addition to the above textual features, the visual content of images is represented using various low-level global image features and several derived features intended to capture high-level semantic content. Low-level Global Features We represents the spatial structure and global shape/edge features of images with the Color Layout Descriptor (CLD) and Edge Histogram Descriptor (EHD) of MPEG-7 [1]. CLD is extracted to form the feature vector f_{old} and EHD is extracted to form f^{ehd} . Additionally, we extract the Color and Edge Directivity Descriptor (CEDD) and Fuzzy Color and Texture Histogram (FCTH) using the Lucene image retrieval (LIRE) library. CEDD incorporates color and texture information into f^{cedd} and FCTH uses the high frequency bands of the Haar wavelet transform to form f^{fcth} .

c) "Bag of Concepts" Feature

In a heterogeneous medical image collection, it is possible to identify specific local patches in images that are perceptually and/or semantically distinguishable, such as homogeneous texture patterns in gray-level radiological images, differential color and texture structures in microscopic pathology and dermoscopic images. The variation in the local patches can be effectively modeled as "visual concepts" [8] by using supervised learning-based classification techniques, such as Support Vector Machines (SVMs). For concept model generation, we utilize a multi-class SVM composed of bi- nary SVM classifiers combined using the one-against-one strategy [3]. To train the SVM, a set of L labels are assigned as $C = \{C_1, \dots, C_j\}$ where each $C_i \in C$ characterizes a visual concept. The training set consists of local patches generated by a fixed-partition and represented by a combination of color and texture moment-based features. The input to the system is the feature vectors for patches along with their manually assigned concept labels. Concept labels are assigned by fixed partitioning each image I_j into l regions as $\{X_{1j}, \dots, X_{kj}, \dots, X_{lj}\}$ where each $X_{kj} \in \mathbb{R}^d$ is a combined color and texture feature vector. For each X_{kj} , its category cm is determined by the prediction of the multi-class SVM. Hence, instead of the low-level feature-based representation, an entire image is represented as

a two-dimensional index linked to visual concepts. Based on this encoding scheme, an image I_j is represented as a vector of concepts f concept.

$$j = [w_{1j}, w_{2j}, \dots, w_{lj}]^T$$

Where each w_{ij} denotes the "tf-idf" weight of a concept ci ; $1 \leq i \leq L$ in image I_j , depending on its information content.

d) "Bag of Keypoints" Feature

Robust and invariant image features are also extracted that are commonly termed affine region detectors [6]. These regions simply refer to a set of pixels or interest points, which are invariant to affine transformations as well as occlusion, lighting, and intra-class variations. We use the Harris-affine detector to locate interest points [5] as a large number of overlapping regions. We then associate with each interest point a vector descriptor invariant to viewpoint changes and, to some extent, illumination changes computed from the intensity pattern within the point. We use a local descriptor developed by Lowe [4] based on the Scale-Invariant Feature Transform (SIFT), to describe the information in a set of scale-invariant coordinates. The SIFT descriptor is chosen to be invariant to viewpoint changes and, to some extent, illumination changes, and to discriminate between the regions. The above features are vector quantized by a self-organizing map (SOM)-based clustering. Finally, images are represented by a bag of these quantized features (i.e., a bag of keypoints). Hence, the model is applied to images by using a visual analogue of the bag of words model used in text retrieval [1].

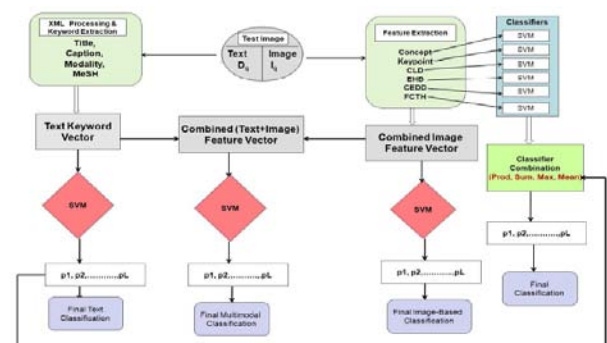


Figure 3 : Process flow diagram of the modality detection approach

V. IDENTIFICATION OF IMAGE-TEXT ASSOCIATION

a) A Similarity-Based Model for Discovering Image-Text Associations

The task of identifying image-text associations can be cast into an information retrieval (IR) problem. Within a web document d containing images and text segments, we treat each image I in d as a query to find

a text segment TS that is most semantically related to I, i.e., TS is most similar to I according to a predefined similarity measure function among all text segments in d. Therefore, we suppose that each image I can be represented as a visual feature vector, denoted as $v^I = (v^I_1, v^I_2, \dots, v^I_n)$, together with a textual feature vector representing the image caption, denoted as $t^I = (t^I_1, t^I_2, \dots, t^I_n)$. For calculating the similarity between an image I and a text segment TS represented by a textual feature vector $t^{TS} = (t^{TS}_1, t^{TS}_2, \dots, t^{TS}_n)$ and hence there is a need to define a similarity measure $\text{sim}_d(<v^I, t^I>, t^{TS})$.

b) Vague-Transformation-Based Cross-Media Similarity Measure

Measuring the similarity between visual and textual features is similar to the task of measuring relevance of documents in the field of multilingual retrieval for selecting documents in one language based on queries expressed in another. For multilingual retrieval, transformations are usually needed for bridging the gap between different representation schemes based on different terminologies [17, 18]. An open problem is that there is usually a basic distinction between the vocabularies of different languages, i.e., word senses may not be organized with words in the same way in different languages. Therefore, an exact mapping from one language to another language may not exist. This is known as the vague problem, which is even more challenging in visual-textual transformation. For individual image regions, they can hardly convey

any meaningful semantics without considering their contexts. For example, a yellow region can be a petal of a flower or can be a part of a flame. On the contrary, words in natural languages usually have a more precise meaning. Some of the models used for Vague-transformation of Image-Text associations are:-

- 1) Statistical Vague Transformation in Multilingual Retrieval.
- 2) Single-Direction Vague Transformation for Cross-Media Information Retrieval.
- 3) Dual-Direction Vague Transformation.
- 4) Vague Transformation with Visual Space Projection.

c) Fusion-ART-Based Cross-Media Similarity Measure

For the vague transformation technique to work on small data sets, we employ an intermediate layer of information categories to map the visual and textual information in an indirect manner. A deficiency of this method is that for training the transformation matrix, additional work on manual categorization of images and text segments is required. In addition, matching visual and textual information based on a small number of information categories may cause a loss of detailed information. Therefore, it would be appealing if we can find a method that learns direct associations between the visual and textual information. Some of the methods used for Fusion- based ART techniques are as follows:

- 1) A Similarity Measure Based on Adaptive Resonance Theory
- 2) Image Annotation Using Fusion ART
- 3) Handling Noisy Text

	Vague Transformation	fusion ART
Approach	Information is translated from one space to another space, so that information from different spaces can be compared.	Information is compared in their respective spaces and the results consolidated based on a multimedia object resonance function.
Learning methodology	Statistic based batch learning.	Incremental competitive learning.
Information encoding	Using transformation matrices to summarize the information based on predefined information categories.	Using self-organizing networks to learn typical categories of multimedia information objects.
Network size	Number of predefined information categories is small as information summarized in the transformation matrices is relatively more compact.	More category nodes are created.
Speed	Around 20 seconds for training; 20 seconds for testing; the training and testing time increases linearly with the number of data samples.	Around 120 seconds for training; 30 seconds for testing; the training and testing time increases exponentially with the number of learnt object templates.
Performance Stability	Unstable	Stable

d) Image and Text Association Presentation

The following description illustrates a process for image and text association:

```

for x to width {
  for y to height {
    blue_diff = blue_of_pixel_in_first_image (x, y) -
    blue_of_pixel_in_second_image(x, y);

```

```

red_diff = red_of_pixel_in_first_image (x, y) -
red_of_pixel_in_second_image(x, y);
green_diff = green_of_pixel_in_first_image (x, y) -
green_of_pixel_in_second_image(x, y);
if (blue_diff < toleration) { blue_diff = 0; }
if (red_diff < toleration) { red_diff = 0; }
if (green_diff < toleration) { green_diff = 0; }

```

```

diff_total = diff_total + ((blue_diff + red_diff +
green_diff) / 3)}
}
Char *SrcImagetmp = (char*) Image Src;
Char *Src2Imagetmp = (char*) Image Src2;
Const char *Imageend = (char*) Image
Src+width*height*3;

While (SrcImagetmp < Imageend)
{blue_diff = *SrcImagetmp++ - *Src2Imagetmp++;
red_diff = *SrcImagetmp++ - *Src2Imagetmp++;
green_diff = *SrcImagetmp++ - *Src2Imagetmp++;

If (blue_diff < toleration) {blue_diff = 0 ;}
If (red_diff < toleration) {red_diff = 0 ;}
If (green_diff < toleration) {green_diff = 0 ;}

diff_total = diff_total + ((blue_diff + red_diff + green_diff) /
3)}

```

VI. EDGE MODEL FOR IMAGE AND TEXT ASSOCIATION

The purpose of detecting sharp changes in image brightness is to capture important events and changes in properties of the world. It can be shown that under rather general assumptions for an image formation model, discontinuities in image brightness are likely to correspond to:

- Discontinuities in depth,
- Discontinuities in surface orientation,
- Changes in material properties and
- Variations in scene illumination.

In the ideal case, the result of applying an edge detector to an image may lead to a set of connected curves that indicate the boundaries of objects, the boundaries of surface markings as well as curves that correspond to discontinuities in surface orientation. Thus, applying an edge detection algorithm to an image may significantly reduce the amount of data to be processed and may therefore filter out information that may be regarded as less relevant, while preserving the important structural properties of an image. If the edge detection step is successful, the subsequent task of interpreting the information contents in the original image may therefore be substantially simplified. However, it is not always possible to obtain such ideal edges from real life images of moderate complexity. Edges extracted from non-trivial images are often hampered by *fragmentation*, meaning that the edge curves are not connected, missing edge segments as well as *false edges* not corresponding to interesting phenomena in the image – thus complicating the subsequent task of interpreting the image data.

Edge detection is one of the fundamental steps in image processing, image analysis, image pattern recognition, and computer vision techniques. During recent years, however, substantial (and successful) research has also been made on computer vision

methods that do not explicitly rely on edge detection as a pre-processing step.

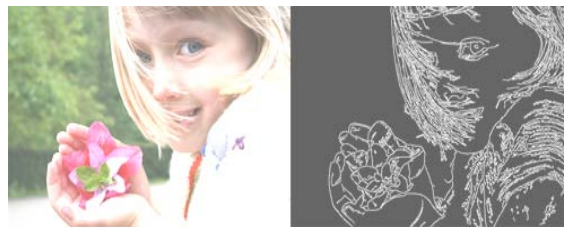


Figure 4 : An image depicting Edge properties

a) Edge Properties

The edges extracted from a two-dimensional image of a three-dimensional scene can be classified as either viewpoint dependent or viewpoint independent. A *viewpoint independent edge* typically reflects inherent properties of the three-dimensional objects, such as surface markings and surface shape. A *viewpoint dependent edge* may change as the viewpoint changes, and typically reflects the geometry of the scene, such as objects occluding one another.

A typical edge might for instance be the border between a block of red color and a block of yellow. In contrast a line (as can be extracted by a ridge detector) can be a small number of pixels of a different color on an otherwise unchanging background. For a line, there may therefore usually be one edge on each side of the line.

b) A simple edge model

Although certain literature has considered the detection of ideal step edges, the edges obtained from natural images are usually not at all ideal step edges. Instead they are normally affected by one or several of the following effects:

- Focal blur caused by a finite depth-of-field and finite point spread function.
- Penumbral blur caused by shadows created by light sources of non-zero radius.
- Shading at a smooth object.

A number of researchers have used a Gaussian smoothed step edge (an error function) as the simplest extension of the ideal step edge model for modeling the effects of edge blur in practical applications. Thus, a one-dimensional image f which has exactly one edge placed at $x = 0$ may be modeled as:

$$f(x) = \frac{I_r - I_l}{2} \left(\operatorname{erf} \left(\frac{x}{\sqrt{2}\sigma} \right) + 1 \right) + I_l.$$

At the left side of the edge, the intensity is $I_l = \lim_{x \rightarrow -\infty} f(x)$, and right of the edge it is $I_r = \lim_{x \rightarrow \infty} f(x)$. The scale parameter σ is called the blur scale of the edge.

VII. TEXT PROCESSING AND RETRIEVAL

To search a body of information (such as a collection of images or citations to the biomedical literature) for objects (scientific articles, images, case descriptions, etc.) relevant to a search query, take the following steps: (1) automatically represent the body of information in the structured form required by our search engine (ITSE) (2) formally represent the information need submitted by a user or inferred from a patient's case using the Evidence Based Medicine framework of a well-formed question; and (3) translate the formal representation of the information need to a search query using the search engine query language. The body of information for our first research initiative consists of full-text scientific publications, whereas in our second initiative we process any type of free text associated with images stored in databases, or find image-related text in an article or case description that contains images.

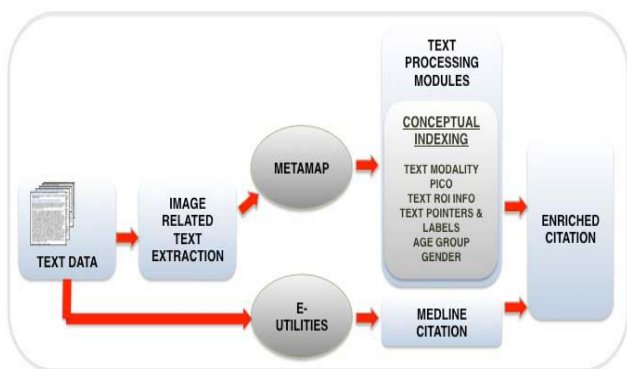


Figure 5 : Text Processing and Reterival

a) *Extracting image-related text: caption segmentation and mention extraction*

The first step in generation of an enriched citation is finding image captions and mentions. In most cases, finding image captions amounts to parsing the corresponding XML tags in such documents as PubMed Central articles or using regular expressions.

b) *Understanding image description: pointers and ROI*

The extracted image-related text is further processed to identify Image Regions of Interest (ROI). ROIs are commonly described in the image caption and indicated by an overlay that facilitates locating the ROI. This is especially true for hard to interpret scientific images such as radiology images. ROIs are also described in terms of location within the image, or by the presence of a particular color.

c) *Image representation (conceptual indexing)*

To provide a structured summary of the salient image content akin to that of the MeSH indexing of the biomedical articles, we explored MetaMap-based extraction of the salient indexing terms from the image-related text[9]. Studies show that image captions provide up to 80% of indexing terms (as judged by

physicians trained in medical informatics), but additional filtering is needed for acceptable precision.

d) *Generating structured documents (enriched citations) for retrieval*

Working with a collection of structured data in XML format provides ready access to document fields, such as title, abstract, conditions, treatment, keywords, etc. Access to document structure supports differential weighting of the occurrences of search terms in different document fields. Structured documents also support faceted search and document and query frame unification (if queries are represented using the same structure). Although the Essie search engine presently supports only the differential weighting of the search terms, creating the structured documents is worthwhile, as we can adjust the field weights for different tasks.

e) *Automatic query generation*

The extraction and grouping approaches are based on the four components of a well-formed clinical question: **P**atient/**P**roblem, **I**ntervention, **C**omparison, and **O**utcome (PICO). To construct the PICO frames, Essie, MetaMap or eXtractor (CTX) for clinical narrative to map the text to the UMLS Metathesaurus and extract concepts relating to problems, interventions (drugs, therapeutic and diagnostic procedures), and anatomy are used.

f) *Essie*

The Essie search engine, developed and used at NLM, features a number of strategies aimed at alleviating the need for sophisticated user queries[10]. These strategies include a fine-grained tokenization algorithm that preserves punctuation, and phrase searching based on the user's query. Essie is particularly well-suited for information retrieval tasks in the medical domain since it performs concept-based indexing, automatically expands query terms using synonymy relationships in the UMLS Metathesaurus, and weights term occurrences according to their document location when computing document scores.

VIII. CONCLUSION

Enriching citations with relevant images can significantly improve literature retrieval for scientific research and clinical decision making. Images are retrieved using image features and visual keywords developed to describe their content. The visual keywords are used to find similar images, followed by IR techniques to improve the relevance of the visually similar images retrieved. Results show no benefit in combining textual and visual features for the retrieval tasks; modality detection is improved when utilizing both text- and content-based approaches. There are two distinct methods for learning and extracting associations between images and texts from multimedia web documents. The vague-transformation based method

utilizes an intermediate layer of information categories for capturing indirect image-text associations. The experimental results suggest that vague method is able to efficiently learn image-text associations from a small training data set. Also, it performs significantly better than the baseline performance provided by a typical image annotation model.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Tao Jiang and Ah-Hwee Tan, *Senior Member, IEEE* "Learning Image-Text Associations" *Manuscript received May 24, 2007; revised January 17, 2008; accepted June 25, 2008.*
2. H K Sawant AWANT, ² Dipali Kadam "Context Based Approaches to Learn Text and Image Association and Processing Semantics" *Journal of IKRIT* ISSN: 0975-6698 Nov. 10 Oct 11 Volume-01 Issue-02 Page 82-88.
3. Chang, S.F., Sikora, T., Puri, A.: Overview of the MPEG-7 standard. *IEEE Transactions on Circuits and Systems for Video Technology* 11(6), 688{695 (2001).
4. Demner - Fushman, D., Lin, J.: Answering clinical questions with knowledge-based and statistical techniques. *Computational Linguistics* 33(1), 63{103 (Mar 2007).
5. Hastie, T., Tibshirani, R.: Classification by pairwise coupling. *The Annals of Statistics* 26(2), 451{471 (1998).
6. Lowe, D.G.: Distinctive image features from scale-invariant key points. *International Journal of Computer Vision* 60(2), 91{110 (2004).
7. Mikolajczyk, K., Schmid, C.: An affine invariant interest point detector. In: *Proceedings of the European Conference on Computer Vision*. pp. 128 142 (2002).
8. Mikolajczyk, K., Tuytelaars, T., Schmid, C., Zisserman, A., Matas, J., Schaals, F., Kadir, T., Gool, L.V.: A comparison of affine region detectors. *International Journal of Computer Vision* 65(1{2), 43{72 (2005).
9. Müller, H., Kalpathy-Cramer, J., Eggel, I., Bedrick, S., Charles E. Kahn, Jr., C.E., Hersh, W.: Overview of the clef 2010 medical image retrieval track. In: *Working Notes of CLEF 2010* (2010).
10. Rahman, M.M., Antani, S., Thoma, G.: A medical image retrieval framework in correlation enhanced visual concept feature space. In: *Proceedings of the 22nd IEEE International Symposium on Computer-Based Medical Systems* (2009).
11. Aronson AR. Effective mapping of biomedical text to the UMLS Metathesaurus: the MetaMap program. *Proc AMIA Symp.* 2001:17-21.
12. McCray AT, Ide NC, Loane RR, Tse T. Strategies for supporting consumer health information seeking. *Stud Health Technol Inform.* 2004; 107(Pt 2):1152-6.
13. D. Radev, "A Common Theory of Information Theory from Multiple Text Sources, Step One: Cross-Document Structure," *Proc. First ACL SIG dial Workshop Discourse and Dialogue*, 2000.
14. R. Barzilay, "Information Fusion for Multi-document Summarization: Paraphrasing and Generation," PhD dissertation, 2003.
15. H. Alani, S. Kim, D.E. Millard, M.J. Weal, W. Hall, P.H. Lewis, and N.R. Shadbolt, "Automatic Ontology-Based Knowledge Extraction from Web Documents," *IEEE Intelligent Systems*, vol. 18, no. 1, pp. 14-21, 2003.
16. S.-F. Chang, R. Manmatha, and T.-S. Chua, "Combining Text and Audio-Visual Features in Video Indexing," *Proc. IEEE Int'l Conf. Acoustics, Speech, and Signal Processing (ICASSP '05)*, pp. 1005-1008, 2005.
17. J. Geurts, S. Bocconi, J. van Ossenbruggen, and L. Hardman, "Towards Ontology-Driven Discourse: From Semantic Graphs to Multimedia Presentations," *Proc. Second Int'l Semantic Web Conf. (ISWC '03)*, pp. 597-612, 2003.
18. Q. Ding, Q. Ding, and W. Perrizo, "Association Rule Mining on Remotely Sensed Images Using P-Trees," *Proc. Sixth Pacific-Asia Conf. Advances in Knowledge Discovery and Data Mining (PAKDD '02)*, pp. 66-79, 2002.
19. I.K. Sethi, I.L. Coman, and D. Stan, "Mining Association Rules between Low-Level Image Features and High-Level Concepts," *Proc. SPIE*, vol. 4384, no. 1, pp. 279-290? PSI/4384/279/1, 2001.
20. Learning Image-Text Associations Tao Jiang and Ah-Hwee Tan, *Senior Member, IEEE* *IEEE TRANSACTIONS ON KNOWLEDGE AND DATA ENGINEERING*, VOL. 21, NO. 2, FEBRUARY 2009.

This page is intentionally left blank



Edge Identification During Fire Environment for Robot

By Md: Rabiul Hasan

BGC Trust University Bangladesh

Abstract - It is very important for us to rescue from fatal accidents caused by fire. Recently in Bangladesh more than two hundred garment workers have deceased at Tajrin Fashion Industry. In this work we have performed the edge detection for Robot on the eve of edge identification to save the worker while they will be locked at emergency situations where human interaction will be failed. We have trained the system such a way that a Robot can easily learn the situations. We have used Automated Brained Learning (ABL) for Robot to detect the objects. Our work ensures only the edge detection. Sobel operator and masking is used in this process. To accomplish the work RGB color model and other color model such as YMC color model is analyzed to ensure the better result. We have noticed that RGB color model is better for our ABL process. Besides, YMC color model also generate good result while the fire is over smoked.

Keywords : *automated brained learning, robot, RGB color model, YMC color model, sobel operator, masking.*

GJCST-F Classification: *1.5.m*



Strictly as per the compliance and regulations of:



Edge Identification During Fire Environment for Robot

Md: Rabiul Hasan

Abstract - It is very important for us to rescue from fatal accidents caused by fire. Recently in Bangladesh more than two hundred garment workers have deceased at Tajrin Fashion Industry. In this work we have performed the edge detection for Robot on the eve of edge identification to save the worker while they will be locked at emergency situations where human interaction will be failed. We have trained the system such a way that a Robot can easily learn the situations. We have used Automated Brained Learning (ABL) for Robot to detect the objects. Our work ensures only the edge detection. Sobel operator and masking is used in this process. To accomplish the work RGB color model and other color model such as YMC color model is analyzed to ensure the better result. We have noticed that RGB color model is better for our ABL process. Besides, YMC color model also generate good result while the fire is over smoked.

Keywords : automated brained learning, robot, RGB color model, YMC color model, sobel operator, masking.

I. INTRODUCTION

Accidents under fire are very dangerous for human being as well as all other animals. Some time it becomes very difficult for respective personnel to save from fired palace due to unavoidable situations. Here we have concentrate to train the Robot for safety purposes. We know that an image is an artifact that depicts or records visual perception. Images are generally represented as two-dimensional arrays (i.e., matrices), in which each element of the matrix corresponds to a single pixel in the displayed image. Pixel is derived from picture element and usually denotes a single dot on a computer display. For example, an image composed of 200 rows and 300 columns of different colored dots would be represented as a 200-by-300 matrix. Some images, such as true color images, require a three-dimensional array, where the first plane in the third dimension represents the red pixel intensities, the second plane represents the green pixel intensities, and the third plane represents the blue pixel intensities.

Segmentation is the technique of partitioning a image into multiple parts [1-3]. Partitions are different components in image which have the similar properties. The outcome of image segmentation is a set of parts that collectively cover the entire image, or a set of

contours extracted from the image. All of the pixels in an area are similar with respect to some properties or computed property, such as color, intensity, or texture. Adjacent regions are significantly different with respect to the same characteristics. Edge detection is one of the most frequently used techniques in digital image processing [3].

The surroundings of image surfaces in a scene often lead to directed localized changes in intensity of an image, called edges. This examination gathered with a commonly held belief that edge detection is the first step in image segmentation, has fueled a long search for a good edge detection algorithm to use in image processing [4]. This search has designed a principal area of research in machine environment and has led to a steady stream of edge detection algorithms for Robot. Edge detection of an image helps to reduces significantly the amount of data and filters out information that may be regarded as less relevant, preserving the important structural properties of an image. For an image designing model, irregularities in image brightness are likely to correspond to a) Discontinuities in height b) Irregularities in surface orientation c) Changes in Particles properties d) changes in scene illumination e) Grayness ambiguity f) Vague knowledge.

II. TYPES OF DERIVATIVES

In this research activity we have used both first order and second order derivatives for edge detections. The measurements of the Gradients in this work are performed base on the equation from 1 to 5 as below. All the equations her are first order derivatives.

$$G(x,y) = \frac{\partial F(x,y)}{\partial x} \cos\theta + \frac{\partial F(x,y)}{\partial y} \sin\theta \quad (1)$$

$$G(j,k) = \sqrt{[G_R(j,k)]^2 + [G_G(j,k)]^2} \quad (2)$$

$$G_R(j,k) = F(j,k) - F(j,k-1) \quad (3)$$

$$G_G(j,k) = F(j,k) - F(j+1,k) \quad (4)$$

$$G_I(j,k) = F(j,k) - F(j+1,k+1) \quad (5)$$

Author : Computer Science and Engineering, BGC Trust University Bangladesh. Roll: 1014023.
E-mail : rabiulhasancsebangladesh@gmail.com

$$G_2(j, k) = F(j, k+1) - F(j+1, k) \quad (6)$$

For second order derivatives we have used following equations. Such as

$$G(j,k)=MAX[|G1(j,k)|,.....|Gm(j,k)|,.....|GM(j,k)|]$$

$$Gm(j, k) = F(j, k) \otimes Hm(j,k)$$

7

$$G(j, k) = \text{Max}[|5 Si-3Ti|]$$

$$i=0$$

$$S_i = A_i + A_{i+1} + A_{i+2}$$

$$S_i = A_{i+3} + A_{i+4} + A_{i+5} + A_{i+6} + A_{i+7}$$

$$\nabla^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2}$$

$$G(x, y) = -\nabla^2 \{F(x, y)\}$$

$$G(j,k)=[F(j,k)-F(j,k-1)]-[F(j,k+1)-F(j,k)]+[F(j,k)-F(j+1,k)]-[F(j-1,k)-F(j,k)]$$

III. DATA COLLECTIONS

We have collected data by creating real fire at both day and night environments. We have clustered the data set by applying the Regression analysis and correlation and variance. Regression analysis helps to better fit the data set properly. Regression analysis, in general sense, means the calculation or measurement of the unknown value of one variable from the known value of the other variable. It is one of the most valuable statistical tools which are extensively used in almost all sciences. It is specially used in data classification to study the relationship between two or more variables that are related causally and for the estimation of demand and supply graphs, cost functions, production and consumption functions and so on. The regression equation for our work is as follows.

$Y^* = a + bx$, where constant value of a and b are.

$$b = \frac{n \sum xy - (\sum x)(\sum y)}{n(\sum x^2) - (\sum x)^2}$$

$$a = \frac{\sum y - b \sum x}{n}$$

IV. AUTOMATED BRAINED LEARNING

Automated Brained Learning represents a brain analogy for information processing. These models are

biologically exhilarated rather than clear-cut clone of how the brain actually functions. Learning by Brained ideas are usually implemented as system simulations of the massively parallel processes that involve processing elements interconnected in network architecture.

$$y_k^1 = \frac{1}{1 + e^{-w^{1kT}x - a_k^1}}, k = 1, 2, 3$$

$$y^1 = (y_1^1, y_2^1, y_3^1)^T$$

$$y_k^2 = \frac{1}{1 + e^{-w^{2kT}y^1 - a_k^2}}, k = 1, 2$$

$$y^2 = (y_1^2, y_2^2)^T$$

$$y_{out} = \sum_{k=1}^2 w_k^3 y_k^2 = w^{3T} y^2$$

$$\Delta w_i^j = -c \cdot \frac{\partial E}{\partial w_i^j}(W)$$

$$w_i^{j,new} = w_i^j + \Delta w_i^j$$

V. IMPLEMENTATION AND RESULT

Based on our Brained Learning we have got following result shown at figure 1 at various moments.



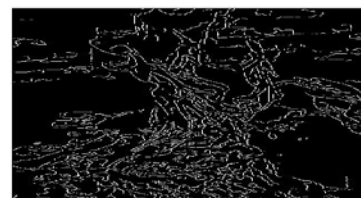
(a)



(b)



(c)



(d)



(e)



(f)



(g)



(h)



(i)

Figure 1 : From a to i we see that the edge detection is implemented with ABL process and proper derivatives

VI. CONCLUSION

We have noticed that our method is working very accurate when the fire intensity is high and it provides poor result at the time of smoked conditions. In future we will try to design algorithm for smoked removal detections.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Orlando J. Tobias and Rui Seara, "Image Segmentation by Histogram Thresholding Using Fuzzy Sets", *IEEE Transactions on Image Processing*, Vol.11, No.12, 2004, pp. 1457-1465.
2. M. Abdulghafour, "Image segmentation using Fuzzy logic and genetic algorithms", *Journal of WSCG*, vol. 11, no.1, 2003.

3. N. Senthilkumaran and R. Rajesh, "Edge Detection Techniques for Image Segmentation-A Survey", *Proceedings of the International Conference on Managing Next Generation Software Applications (MNGSA-08)*, 2008, pp.749-760.
4. N. Senthilkumaran and R. Rajesh, "A Study on Edge Detection Methods for Image Segmentation", *Proceedings of the International Conference on Mathematics and Computer Science (ICMCS-2009)*, 2009, Vol.I, pp.255-259.
5. Chan, R.H., Ho, C.W., and Nikolova, M., 2005. Salt and Pepper Noise Removal by Median Type Noise Detectors and Detail-Preserving Regularization. *IEEE Transactions on Image Processing*, vol. 14, no.10, pp. 1479-1485.
6. Brownrigg, D.R.K., 1984. The weighted median filter. *Communication, ACM*, vol. 27, no.8, pp. 807-818.



This page is intentionally left blank



Performance Analysis of Gray Scale and Color Iris with Multidomain Feature Normalization and Dimensionality Reduction

By V. V. Satyanarayana Tallapragada & Dr. E. G. Rajan

Pentagram Research Center, Hyderabad

Abstract - Iris recognition techniques have come a long way since the preliminary proposal by Daugman. There are several techniques which mainly focus on improving the performance of IRIS recognition system either with the help of improved classifier or the feature set. However we have proved that the performance of IRIS recognition to a large extent depends upon the inter dependability and separability of the feature vectors from different classes in the feature plane. If such dependency is identified and dimensions or the features for which most of the classes represent similar and inseparable values. This finding lead us to investigate the dimension reductionality on the feature vectors.

Most of the Iris recognition technique still relies on Gabor filter. But modern IRIS recognition sensors present a great deal of details about the IRIS part and the resultant images are color image. Hence there is a need to analyze the behavior of dimensionality reduction techniques for both gray scale conventional IRIS recognition technique as well as the one applied on color images.

Keywords : *IRIS, gabor filter, features, dimensionality reduction.*

GJCST-F Classification: *1.4.5*



PERFORMANCE ANALYSIS OF GRAY SCALE AND COLOR IRIS WITH MULTIDOMAIN FEATURE NORMALIZATION AND DIMENSIONALITY REDUCTION

Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

Performance Analysis of Gray Scale and Color Iris with Multidomain Feature Normalization and Dimensionality Reduction

V. V. Satyanarayana Tallapragada^a & Dr. E. G. Rajan^o

Abstract - Iris recognition techniques have come a long way since the preliminary proposal by Daugman. There are several techniques which mainly focus on improving the performance of IRIS recognition system either with the help of improved classifier or the feature set. However we have proved that the performance of IRIS recognition to a large extent depends upon the inter dependability and separability of the feature vectors from different classes in the feature plane. If such dependency is identified and dimensions or the features for which most of the classes represent similar and inseparable values. This finding lead us to investigate the dimension reductionality on the feature vectors.

Most of the Iris recognition technique still relies on Gabor filter. But modern IRIS recognition sensors present a great deal of details about the IRIS part and the resultant images are color image. Hence there is a need to analyze the behavior of dimensionality reduction techniques for both gray scale conventional IRIS recognition technique as well as the one applied on color images.

In this work we apply the dimensionality reduction technique on multi-domain features extracted from the images for set of color and gray scale IRIS dataset. The result shows that the dimensionality reduction on color images improves the performance of the classifier by .8% and the performance of the gray scale image database classifier is improved by .3%.

Keywords : IRIS, gabor filter, features, dimensionality reduction.

I. INTRODUCTION

The features of a pattern are represented in different scales and dimensionality. In a huge dataset of pattern values it becomes difficult to identify the underneath correlation amongst the values and in turn amongst the objects. For example, in a set of color feature matrix of the size that of the size of the image, if we extract min, max, mean, deviation then we can classify another matrix to be similar to this or not.

Such statistics are known as low dimensional data because these are extracted from direct functions of data itself.

Such low dimensional data are densely correlated in the inter-class spacing where each class represents a specific pattern or set of patterns distinguishable from one another. The spacing or the

decision boundary amongst these data or data points determines whether a given pattern can be classified correctly as one of these patterns or not. Because of high density spacing of these data it is difficult to define a clear decision boundary. Hence a common approach is to represent the behavior of the data through more dimensions. For example, representing the just mentioned dataset through mean, standard deviation, kurtosis, and moments would mean representing the data through more dimensions. But n-dimensional data has a distinct problem of defining a relationship amongst the dimensions. Therefore a new approach is developed to represent such data through high dimensions where dimensions are mapped through some distinct defined functions. Such kind of data representation is commonly known as high dimensional data. It is believed that higher the representing dimensions of the data, more unique will be the pattern represented by the dimensions together. For example, only a limited behavior is extracted from a time scale analysis of a one dimensional signal. More information is extracted when the same signal is mapped to frequency plane, Laplacian plane, Phase plane and so on. The same thing stands true even for images. An image is either two-dimensional or three-dimensional spatial information. But x, y and z dimensions of the image pixels can reveal only visual appearance of the image and not the underneath pattern. An image is recognized through various techniques based on the type of pattern it possesses. For example, shape [12,13,14,15,16], texture [17,18,19,20,21,22,23], color [7,8,9,10,11], statistics [24] and so on. There are numerous different methodologies proposed in each of these for various kinds of object detection and pattern recognition.

II. PROBLEM FORMATION

In the introduction section we have discussed about the pattern classification problem, the role of data representation and its dimensionality in pattern classification and the most adopted technique for dimensionality reduction using kernel based methods. Kernel based methods also has several drawbacks which needs to be addressed in order to come to a conclusion that an alternative method is in need for a scalable biometric system.

^a Author : Assistant Professor, Dept. of ECE, C.B.I.T., Gandipet, Hyderabad. E-mail : satya.tvv@gmail.com

^o Author : President, Pentagon Research Center, Hyderabad. E-mail : rajaneg@yahoo.co.in

With the increase in database size, the error rate acceptance decreases significantly, making the Biometric system more vulnerable for false alarm and templates are generated by quantizing and linearization of a particular iris feature.

IRIS features are broadly categorized as follows

1. Color Feature
2. Phase based feature
3. Texture feature
4. Shape Features
5. Frequency based Features

Each feature mentioned here have their pros and cons.

Now if we combine various features and generate a template from all features, the template will be much stronger than generated template through a single feature. But the length of the template will be increased. If independent features are stored, memory requirement becomes immense and searching becomes complicated. Further each type of features is also generated through various techniques.

The reason for selecting multidomain features instead of 2d gabor kernel is that gabor kernel assumes the input to be of gray scale which reduces the property of the IRIS images. By using multidomain features, especially the color feature, the dimensional dependency in iris recognition system can be explored to a great details.

Kernel based techniques also relies on some type of dimensional optimization. The main function of a kernel is to transform two dimensional image matrix into one dimensional feature vector. A transformation from matrix to vector enhances the underneath pattern, at the same time reducing the dimension of the image. The optimization process is recursive and depends upon the learning data available. Therefore if the number of classes in the training dataset is very high with very little number of samples per class then the resulting vector of an optimization process from a kernel based technique will fail to produce optimum result.

Kernel based techniques adopt a dimensionality optimization through low dimensional data. This data is nothing but independent features extracted from specific feature extraction techniques like texture feature, shape features and so on. Low number of samples needs more features to be recognized efficiently. Therefore the feature selection technique must incorporate more than one type of methods or techniques to achieve high accuracy. Each of these feature sets have their own data space and can be considered as an independent dimension. Therefore combination of multiple features can be defined as a set of feature points over an n-dimensional feature space. Kernels cannot map high dimensional data into hyper dimensional data. A hyper dimensional data is in itself a multi dimensional vector extracted of n-dimensional

scalar features. This is the core drawback of a kernel based technique in conventional pattern recognition. Considering the fact that a biometric recognition system may need to recognize billions of classes in the future context, the training process of the kernel cannot be considered feasible because of requirement of huge number of training samples. Further the mapping function itself in the case of a kernel technique evolves with change in data and with increase in number of classes. Therefore it can be said that at every scaling juncture in a biometric recognition a kernel needs to relearn and recompute the feature vectors which in turn is practically not feasible for a highly scalable biometric dataset.

Also with the advancement of technology, the pixel density, clarity, color definition of the images has improved considerably. This has opened up new horizons for IRIS recognition systems. Therefore this work specifically focuses on the roadmap that connects the performance of both color as well as gray scale IRIS recognition system.

III. PROPOSED SYSTEM

Considering the limitations presented by kernel based techniques or other techniques which combines different dimensions or even dimension optimization problem, we propose a new approach for IRIS recognition by combining features and further reducing them systematically.

Let us have a hypothetical discussion and consider a matrix of following structure.

Table 1 : Sample Feature Matrix

Phase				Texture				Shape				Frequency			
A	B	C	D	A	B	C	D	A	B	C	D	A	B	C	D

Let us call the structure shown in table 1 as PTSF and A, B, C, D represents four different techniques adopted for each type of features. Now let us consider that each feature is a unique value between -32,000 to 32,000 or of 16 bits. Texture pattern of iris few are bound to be similar, the way shape pattern can be same. Hence no matter however number of bits is considered for generating the iris, there bound to be interclass similarity. Therefore a better solution is to increase the number of features and take a matching decision based on all the features. This method however increases the complexity of computation and also requires more memory for processing and storage. Hence we proposes a multistage dimension reductionality model with feature normalization, for faster optimization of the dimensionalities and improved performance in terms of both time and space computational complexity.

In case of color images, color features are obtained from all three color descriptors (i.e. Red, green and blue component) where as for grasy scale IRIS images, gray scale is combined with local binary pattern

and binary color scale or keep the domain or number of features consistent.

IV. METHODOLOGY

a) IRIS Segmentation

We use Daugman[24] IRIS segmentation technique for MMU Image database and used a custom segmentation for Phoenix database, as Daugman's technique is not suitable for color IRIS segmentation due to dependency of Daugman's technique on 2-D Hough circles.

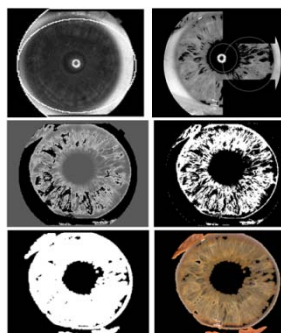


Fig. 2 : Results of Segmentation

The color segmentation used here is explained as below.

1. Scan IRIS Image.
2. Extract the image pixels satisfying following criteria.
3. $I = I < 50$; where I refers to pixel value.
4. Remove the extra noise by eroding the image with a 3×3 structuring element.
5. Fill the holes in the mask.
6. Segment original image with this mask. After segmentation a part of pupil still appears in the image. In order to remove this part, dilate the image with a disk structuring element with $R=3$ and $N=4$.
7. Segment the image with new mask. Step 1 to 7 removes the pupil part efficiently. Now the following steps are adopted to remove the outer boundary or the retina part of the image.
8. Convert the segmented image to gray scale image. Apply histogram equalization. This step enhances the retina part to a brighter color.
9. Remove the brighter retina part with hard thresholding. Pixel values over 200 are made as zero.
10. Convert the image to binary image.
11. Remove independent noise by erosion operator with a square element of size 2×2 .
12. Fill the image by applying dilation with a disk operator of size $R=8$, $N=8$.
13. Re-segment the result of step 7 with the new mask to obtain an image with only the iris part. Sample result of an IRIS segmentation using the above steps is presented in figure 2.

b) Color, Texture, Phase and Shape Feature Extraction

i. Color Feature Extraction

- (i) As each iris comprise of a uniform color distribution we consider the color variance rather than the mean color values.
- (ii) Following features are extracted out of red, green and blue components.
 - a) Standard deviation (Dimension 1)
 - b) Variance (Dimension 2)
 - c) Kurtosis (Dimension 3)
 - d) Mean of non NAN Co-Variance (Dimension 4)
 - e) Statistical moments of order III (Dimension 5) and IV (Dimension 6)

These six sample features extracted from an iris image are shown in table 2.

Table II : Values of the 6 Features as Extracted from Results of Fig 2

35112.4427978787	5574.25561362929	3511.44954063009
1.75451300297860	1.41677277950672	1.64405543018030

The technique used for local binary pattern is as follows.

- 1) Define a window size W
- 2) Scan every pixel in an image, extract its $W \times W$ neighbourhood.
- 3) Check if a Neighbour pixel color $>$ Center, if so put 1 in the matrix else 0
- 4) Thus we get an array of '1's and '0's'. Convert this to binary and subsequently binary needs to be converted to decimal and must replace the center pixel. Remember bigger the value of W , larger will be number. For example for $W=4$, you will get a binary number of 16 bits. But gray scale image can contain only 8 bit colors. Therefore once entire image's LBP is extracted, it is to be normalized.
- 5) Normalization is performed as $(\text{Image containing local binary pattern}) * 255 / (\text{maximum value in the image})$

Important thing is if you take $W=1$, resultant image will be fine edge detected image. So you can alternatively use this theory to extract edges from images.

One might argue the point of using LBP as color feature as mainly binary pattern represents the texture of a surface. As the gray scale images are low in separable properties in the color domain, representing local binary pattern as colors improves the separability of the color features for gray scale images. It is quite clear that dimension optimization through Principal Components is not feasible where the data domains are not normalizable.

ii. Shape Feature Extraction

Shapes in an IRIS pattern [21,22] are nothing but the texture lines obtained in the image. Edges describe this pattern in the best manner. The main problem with considering edge information of an IRIS is

that the edge information of IRIS depends upon the venation pattern in IRIS and not its edge from boundary like any other objects. This venation pattern further depends upon the quality of scanning and preprocessing. Therefore we transform the image to polar coordinate system. As it is apparent from figure 3 that average IRIS image edges belong from zero to ninety degrees.

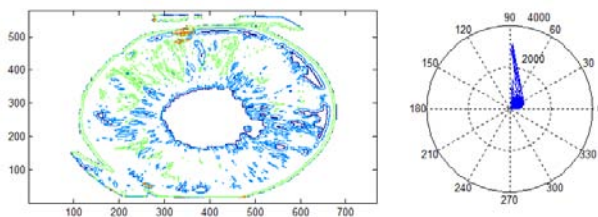


Fig. 3 : Contour of Segmented IRIS and its Polar Transformation Data

We extract the arc length from rho and the metrics corresponding to polar transform using $L = \rho \cdot \theta$

Skewness(Dimension 7), Mean(Dimension 8) and standard deviation(Dimension 9) are extracted as shape features. Further the real part of 10th order Zernike moments [16] from the image are extracted (Dimension 10 to 15) and are combined with polar transform features to obtain shape features. The values extracted are shown in table 3.

Table III : Sample Shape Features Combining Zernike Moments and Features from Polar Component of Image Contour

7.72E+05	-4.64959.537	4.27E+05	-1.86E+05	1.04E+05
2.22E+05		1.91E+04	1.13E+04	8.90E-01

For gray scale MMU database, Variance of the Gray scale components of the segmented IRIS part is considered.

iii. Texture Feature Extraction

Texture pattern in an IRIS is a common type of feature. Each IRIS represents a unique texture pattern which can be represented by various scaling of the images or co-occurrence matrix. Co-occurrence matrix in the gray level posses several properties like energy, entropy, average intensity which are considered as good descriptors of the IRIS pattern.

iv. Grey Level Co-Occurrence Matrix (GLCM)

A cooccurrence is defined as probability of repeating the same color value in neighborhood pixels in any conventional scan like horizontal, diagonal or vertical. Subsequently gray level cooccurrence matrix can be defined as the probability of occurrence of a gray level color with corresponding to all other colors

and the self in any or all mentioned scanning directions. As this is a probability matrix the feature of the entire matrix can be represented by probabilistic statistical measurement or analysis like entropy(mutual information), energy(maximum information contained), intensity (probabilistic intensity distribution statistics), mean and standard deviation (1st order statistical moments), contrast(the brightness factor), autocorrelation(probability of same probabilistic dependency amongst different colors), homogeneity (probability of uniformity of the color distribution), dissimilarity(probability of variance amongst different colors), inertia (the central moment representing the overall distribution of the matrix) and so on. We have considered only 7 specific properties in this work and are Energy, Dissimilarity, Homogeneity, Correlation, Inertia, Autocorrelation, and Contrast [17].

v. Feature Extraction in Phase and Frequency Domain

Let us assume that the image is scaled to a size of 256 x 256. The magnitude or the phase components will have 256 elements where the most high frequency elements are at the center or around the vertical and the horizontal axis due to the fact that most IRIS images are circular and the central IRIS has most invariability as shown in figure 4.

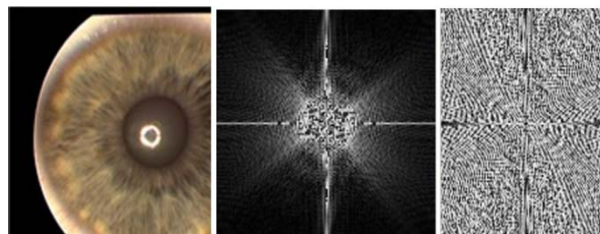


Fig. 4 : Iris Image and its corresponding Phase and Magnitude Plot

The actual IRIS pattern information will be present outside the central block of values. Also that IRIS images will be symmetric around the center. We perform the following steps to reduce 256 x 256 features describing the IRIS pattern.

1. Locate the Maxima on the magnitude plot as shown in figure 5 as a superimposed magnitude plots of seven iris images.

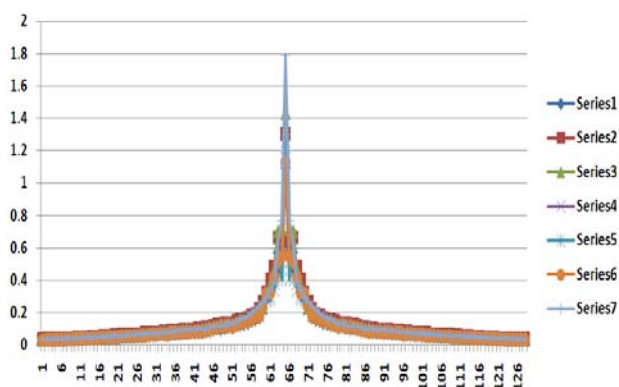


Fig. 5 : Magnitude Plot of 7 IRIS images from the database where FFT is carried out on scaled image of size 128 x 128

2. Leave 15% values on the left and on the right of the central maxima and call them as point1 (p1) and point 2 (p2) respectively.
3. As the IRIS image is symmetric on the center, take two groups: from 1 to p1 and the other group from p2 till the end and call them as group1 and group2.
4. Take the element wise mean of both group1 and group2 to obtain a reduced set of features which will be about 70% of the original features. Thus in a 128 pixel scaled IRIS, total features will be about 90 features. Out of them the lower components will have more invariability than the higher frequency component. Call these feature vectors as set 1.
5. In order to find the significance of the features, extract this pattern from at least 25% of the total training images from database. Extract the principal components from the feature set.
6. The principal component analysis also justifies that the Fourier Magnitude pattern presents maximum information on IRIS images in the low frequency.
7. First find N non Zero components in this curve and retain first N features from the actual set1 as shown in figure 6.

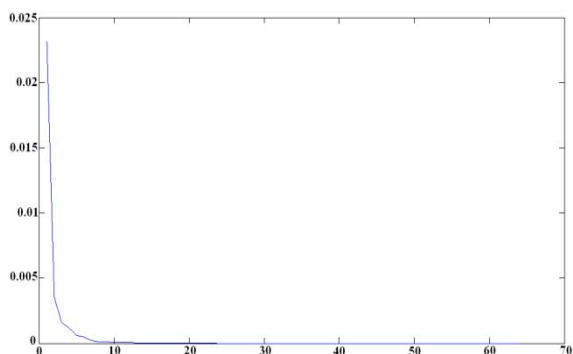


Fig. 6 : Plot of Variance of Principal components of the IRIS magnitude reduced set 1 of scaled IRIS images of size 128 x 128

8. This technique has a unique significance. Fourier features are not the only features which are

attributed here for IRIS recognition purpose. Therefore retaining the principal component deviates the natural feature vectors. Hence we use principal components to eliminate unwanted features from the actual set.

c) Dimension Modeling of the features

Color and texture features are considered to be interrelated and reciprocatory. Hence we consider the color plane (gray scale) and texture plane to be the first two dimensions. The shapes are extracted from the edges which again depend upon the texture values and the changes in the texture. Hence the shape features will be considered as the perpendicular plane over the texture plane.

The features attributed by different feature selection technique have different values which makes it really difficult for dimension reductionality.

1. In order to bring all the feature points in a dimensional plane with quantized values, all the features are self normalized. The normalization is carried out by first obtaining entire set of values of a feature and dividing each feature with the maximum of the domain values. We consider that each color contributes x, y, z component of the independent feature in the color feature space and that the values are between 0 and 1000. The normalized feature set is multiplied by 1000 to range the values between 0 and 1000.

Table IV : Normalized Color Features after Step 1

2.38E+03	3.78E+02	3.78E+02
3.02E+01	2.44E+01	28.31919
5.07E+02	6.88E+02	1.85E+02
3.06E+03	-5.13E+01	-5.45E+00
2.99E+03	6.08E+00	3.64E-01
248.1883	51.68849	248.1883

2. Table 4 depicts a sample color feature set where some values are greater than 1000, and some values are negative. We consider all such values as out of the dimension space and consider the Negative values to be zero and values more than 1000 to be equal to 1000 i.e., on the boundary.

Table V : Color Feature Values after Step 2

Columns 1 through 7						
1000	827	827	144	117	78	207
Columns 8 through 14						
177	377	1000	0	1000	1000	390
		0		0	0	
Columns 15 through 18						
605	182	123	182			

Now even though the overall dimension of the image is huge, the values are uniquely normalized to same domain making it easier for a dimension reduction technique to optimize it.

We adopt two techniques for dimension reduction.

1. Maximum Likelihood Estimation of Intrinsic Dimension
2. Self Organizing MAPS

The basic goal of the entire work is to construct a lattice where all the classes can be uniquely represented. There are over 100 features extracted originally and dimension reductionality of the entire feature set through Self Organizing MAP is near impossible and is of great computational challenge in terms of space and time complexity. Therefore we find out actual "hidden" dimension in the dataset first by applying MLE based Intrinsic Dimensionality Measure. The process of finding out the exact intrinsic dimensionality is an optimization problem which is subjected to continuous assumptions of the probable dimensionality and verifies the belief of the dimensionality. The iterative process of the assumption and the belief measurement is distributed as poisons distribution.

Maximum likelihood estimation of the dimension M of a dataset X can be represented as

$$\hat{m}_k(x) = \left[\frac{1}{k-1} \sum_{j=1}^{k-1} \log \frac{T_k(x)}{T_j(x)} \right]^{-1} \text{ where } k \text{ is the total}$$

number of actual dimensions and T_k is the k^{th} dimension of dataset X and T_j is the j^{th} dimension of the same dataset. It may be noted here that k can also be considered as maximum number of neighbors of any point in the dimensional space rather than the dimensionality itself [27]. Normalized color features over the reduced dimensionality through MLE intrinsic dimensionality reduction are shown in figure 7.

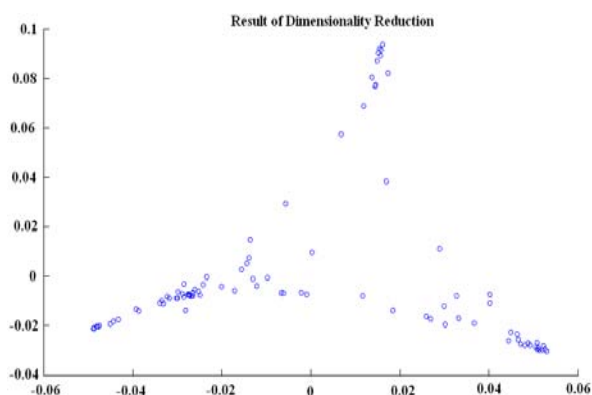


Fig. 7 : Result of MLE

d) Self Organizing Map

The idea of Kohonen for feature reductionality is explained in detail by A.P. Papliński [28]. According to this idea the feature vectors are optimized to a two dimensional SOM plane as shown in figure 8.

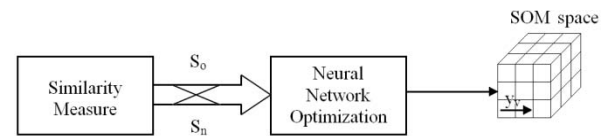


Fig. 8 : Inner block diagram of the optimization process is presented by figure 9

i. Distance Measure

A general block diagram of a Self-Organizing Feature Map along with its parameters is presented in figure 9.

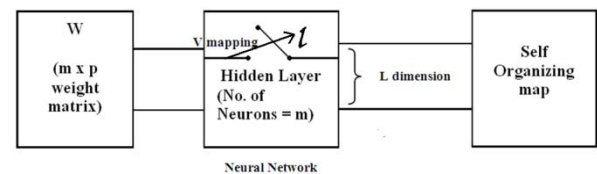


Fig. 9 : Self-Organizing Feature Map in detail [28]

p — Dimensionality of the input space

l — Dimensionality of the neuronal space

m — total number of neurons

W — $m \times p$ Matrix of synaptic weights

V — $m \times l$ Matrix of topological positions of neurons

SOM optimization for set of feature vectors can be summarized as below

a. Initialization

At the beginning the dimensional relationship amongst the data points in the high dimensional space is unknown. Therefore a model is built with a weight matrix W with a randomized sample of input vectors. Learning algorithm is selected.

b. Model Updation

(i) The initial model $M\{m, W\}$ is checked for a winning neuron $k(n)$ and its position $V(k)$ with a condition that the winning neuron is selected in such a manner that its weight difference from neighborhood weight matrix should be larger than the weight differentials of all other competing neuron, where 'n' are the distinct dimensionality of the dataset X . The weight differentials can be measured through any distance formula like Euclidean distance. Finding if a point j is in the neighborhood of 'n' depends upon an exponential distribution centered around the total variance of the dimensional space.

(ii) Once a winning neuron is selected, the weights of all the neurons are updated in reference to the weight of the winning neuron.

c. Neighborhood Shrinking or Reordering

Once a winning neuron is obtained and the neighborhood weights are updated, it is obvious that for the all the feature points in its neighborhood it will always be the winner. Hence shrink the neighbors of the neuron into neuron space itself. The result of the

discussed dimensionality reduction technique based on SOM applied over the feature set of the training iris data is shown in figure 10.

d. Testing

Test features are extracted from the test image with its original dimensionality (in this case reduced by MLE). It is simulated in the optimized neural space to find the position of the winning neuron for the given feature. The classification is done by checking the class whose vectors are in this neuron space.

In the Proposed System a SOM is constructed with final 7 dimensions optimized through MLE intrinsic dimension estimation technique. The SOM comprises of 10 x 64 number of neurons for an experimental evaluation of 64 persons IRIS images.

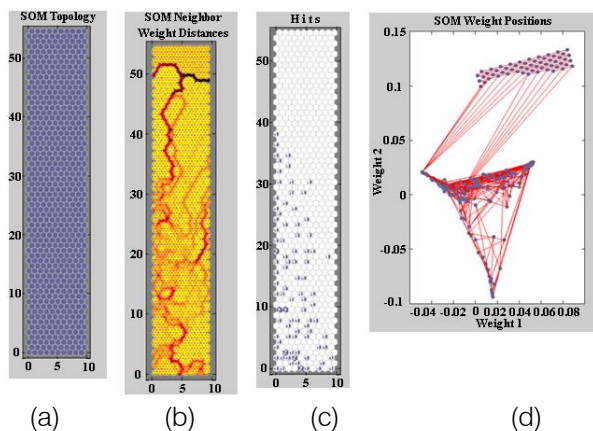


Fig. 10 : (a) The Hexagonal Lattice Structure of the SOM
(b) SOM Neighbour Weight Distances
(c) Number of unique positions affected by neurons after a SOM test
(d) SOM weight position in the reduced dimensional space

V. RESULTS

We use the MMU IRIS database which contains the IRIS images of 100 persons, 5 images of each eye.

DRF for various features and techniques

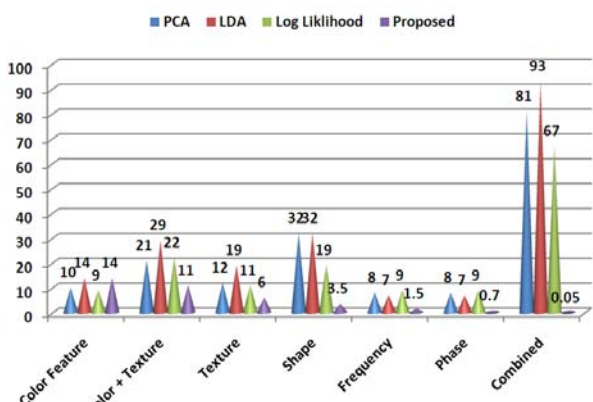


Fig. 11 : Dimension Reduction achieved in different techniques for features and combination of features

We define DRF or dimension Reduction Fraction as $DRF = (\text{Number of Dimension after Reduction}) / \text{Total Dimension before Reduction}$

Figure 11 shows that the proposed technique achieves higher DRF when the Number of Features are high. This is due to the fact that Number of output domain in the proposed technique is fixed to 2. Out of the other techniques, PCA performs better in most of the feature space. The performance of the dimension reduction technique in the combined feature space is not satisfactory. This is purely because PCA or LDA depends upon the high Variance dimension selection. Inter class or feature variance differs a great deal making it difficult for the kernel techniques to achieve high reduction.

Number of features versus DRF for various techniques

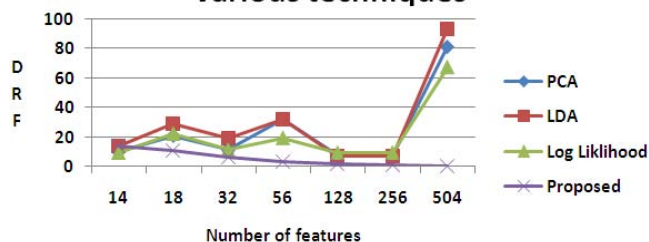


Fig. 12 : Performance of Dimension Reduction techniques for different number of input dimensions

Figure 12 clearly elaborates the fact that various techniques for dimension reduction performs differently for different types of features. They purely vary based on the range of feature values as selected in proposed technique. Therefore clear winner is dimension reduction method.

Reducing the dimension is not the only objective, it must also be proved that reducing the number of dimensions do not reduce the efficiency of classification. Though Kernel based techniques have already shown that reducing the dimensions results in better classification rate due to the fact that the original feature space contains many sparse data and data that are irrelevant for classification decision and when such data or features are considered for classification process, the overall accuracy declines.

Technique versus Efficiency

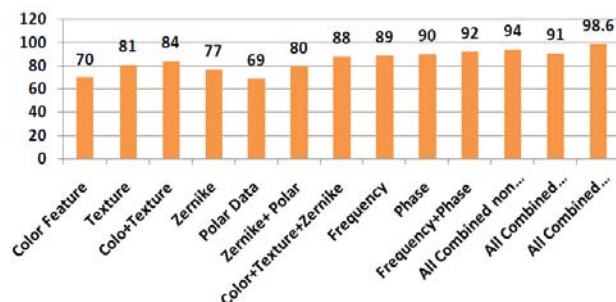


Fig. 13 : Technique Vs Efficiency

VI. CONCLUSION

Figure 13 shows that highest efficiency is achieved by combining the features. By reducing the features through the dimension optimization process do not sacrifice the accuracy to a great deal. In fact it achieves better result over the normalized values. Efficiency after normalization decreases because of loss of information in the course of the process.

Number of features versus RT

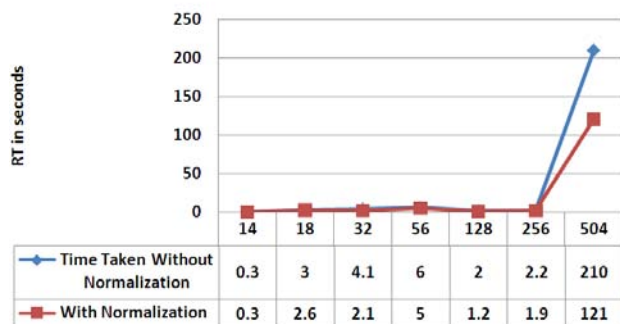


Fig. 14 : Experimental results for reduction time analysis for normalized and non normalized feature vectors

We define reduction time as the time taken for optimization of all features from all the training classes. This is an important experiment where we measure the time taken by the optimization process. Figure 14 clearly proves that with normalization, optimization is faster when number of features is more.

It is clear from the results that overall accuracy of the system is 98.6% when adopted with proper dimension normalization and dimensionality reduction as against 94% of the non normalized features. This is mainly due to sparse vector removal from the feature set.

Figure 15 is obtained by first considering 500 features as standard and obtaining the accuracy of both databases using the method proposed by gabor. Now the accuracy of the proposed system is obtained. The improvement in accuracy of proposed and the standard efficiency is compared.

It can be seen that the color iris images shows a high tendency for improvement due to better separation, especially in the color domain.

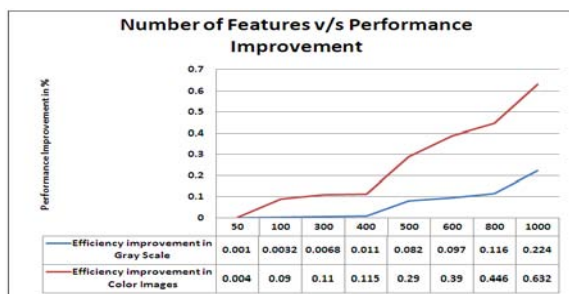


Figure 15 : Performance analysis of the effect of dimensionality reduction over both gray scale and color image dataset

Dimension reduction is relatively a new area of study in IRIS recognition problem. In this paper we have studied the dimensionality reduction problem in great depth and compared the performance of the same with the most extensively used method for IRIS recognition. Results shows that dimensionality reduction improves the overall accuracy of any IRIS recognition system. However multidomain approach outperforms the conventional 2-d gabor based method due to larger set of independent and separable features. As color images present more details than the gray scale images, the improvement is better apparent in color images. The performance of 2-d Gabor filter on each of the independent color domain can be applied as a future work .

REFERENCES RÉFÉRENCES REFERENCIAS

1. Jing Luo, Shuzhong Lin, Ming Lei, Jianyun Ni, "Application of Dimensionality Reduction Analysis to Fingerprint Recognition," *icid*, vol. 2, pp.102-105, 2008 International Symposium on Computational Intelligence and Design, 2008.
2. Wen-Shiung Chen, Chi-An Chuan, Sheng-Wen Shih, Shun-Hsun Chang, "Iris recognition using 2D-LDA + 2D-PCA," *icassp*, pp.869-872, 2009 IEEE International Conference on Acoustics, Speech and Signal Processing, 2009.
3. Tenenbaum, J. B., de Silva, V., & Langford, J. C. (2000). A global geometric framework for nonlinear dimensionality reduction. *Science*, 290, 2319–2323.
4. B. Schölkopf, A. Smola, and K.-R. Müller. Kernel principal component analysis. In B. Schölkopf, C. Burges, and A. Smola, editors, *Advances in Kernel Methods — Support Vector Learning*. MIT Press, Cambridge, MA, 1999b. 327 – 352.
5. S. Lespinats, M. Verleysen, A. Giron, and G. Fertil, "DD-HDS: A method for visualization and exploration of high-dimensional data," *IEEE Trans. Neural Netw.*, vol. 18, no. 5, pp. 1265–1279, Sep. 2007.
6. Zhou Zhiping, Hui Maomao and Sun Ziwen, "An iris recognition method based on 2DWPCA and neural network". *Chinese Control and Decision Conference (CCDC '09)*, 17-19 June 2009, pp. 2357 – 2360. Digital Object Identifier 10.1109/CCDC.2009.5192124.
7. Caitang Sun, Farid Melgani, De Natale Francesco, Chunguang Zhou, Libiao Zhang, Xiaohua Liu, "Incremental Learning Based Color Iris Recognition," *micai*, pp.319-324, 2008 Seventh Mexican International Conference on Artificial Intelligence, 2008.
8. Caitang Sun, Farid Melgani, Chunguang Zhou, De Natale Francesco, Libiao Zhang, Xiangdong Liu, "Semi-Supervised Learning Based Color Iris

- Recognition," *icnc*, vol. 4, pp.242-249, 2008 Fourth International Conference on Natural Computation, 2008.
9. Ryszard S. Choras, "Iris Image Recognition," *cisim*, pp.26-30, 2007 6th International Conference on Computer Information Systems and Industrial Management Applications, 2007.
10. Haishan Chen, Xuan Han, Bochao Hu, "An Image Retrieval Method Based on Spatial Features of Colors," *iscsct*, vol. 1, pp.284-287, 2008 International Symposium on Computer Science and Computational Technology, 2008.
11. Hasan Demirel, Gholamreza Anbarjafari, "Iris Recognition System Using Combined Colour Statistics", *IEEE Symposium on Signal Processing and Information Technology (ISSPIT 2008)*, Sarajevo, Bosnia, Dec. 2008, pp. 175-179.
12. N. Sudha, N. B. Puan, H. Xia, and X. Jiang, "Iris recognition on edge maps," *IET Computer Vision*, vol. 3, no. 1, pp. 1–7, 2009.
13. S. Bhat and M. Savvides. Evaluating active shape models for eyeshape classification. In *IEEE Int. Conf. on Acoustics, Speech, and Signal Processing (ICASSP2008)*, pages 5228–5231, 2008.
14. Ma Zheng, Zheng Tao and Pan Lili, "Classification – Based Algorithms for Iris Outer Edge Location", *IEEE* 2009.
15. Yanyun Zhao, Anni Cai," A Novel Relative Orientation Feature for Shape-Based Object Recognition", *Proceedings of IC-NIDC*, pp. 686-689, 2009.
16. Manjunath G, M. Naresh, V.V.Satyanarayana Tallapragada, "A Novel Shape Recognition Technique by Shape Context and Zernike Moments for Content Based 3-D Object Retrieval System", *International Conference on Systemics, Cybernetics, and Informatics (ICSCI 2011)*, pp. January, 2011.
17. V.V. Satyanarayana Tallapragada, E.G. Rajan, "Iris Recognition Based on Combined Feature of GLCM and Wavelet Transform," *iciic*, pp.205-210, 2010 First International Conference on Integrated Intelligent Computing, 2010.
18. Li Ma, Tieniu Tan, Yunhong Wang, and Dexin Zhang, " Personal Identification Based on Iris Texture Analysis", *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 25, No. 12, pp. 1519-1533, 2003.
19. Fares S. Al-Qunaieer and Lahouari Ghouti, " Color Iris Recognition Using Hypercomplex Gabor Wavelets", *Symposium on Bio-inspired Learning and Intelligent Systems for Security*, pp.18-19, 2009.
20. Li Ma, Tieniu Tan, Yunhong Wang, and Dexin Zhang, " Personal Identification Based on Iris Texture Analysis", *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 25, No. 12, pp. 1519-1533, 2003.
21. S. Liu, H. Yi, L.T. Chia, and D. Rajan. Adaptive hierarchical multi-class svm classifier for texture-based image classification. *Proc. IEEE International Conference on Multimedia and Expo*, page 1-4, 2005.
22. R. Y. F. Ng, Y. H. Tay, and K. M. Mok, "Iris Recognition Algorithms Based on Texture Analysis," *Proc. 3rd International Symposium on Information Technology*, vol. 2, Aug 2008, pp. 904-908.
23. Kamil Grabowski, Mariusz Zubert, Ma_gorzata Napieralska, Andrzej Napieralski, "Uncertainty in Iris Recognition Based on Texture Analysis", *17th International Conference Mixed Design of Integrated Circuits and Systems (MIXDES 2010)*, pp. 587-592, 2010.
24. J. G. Daugman, "High confidence visual recognition of persons by a test of statistical independence," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 15, no. 11, pp. 1148–1161, Nov. 1993.
25. Randy P. Broussard, Lauren R. Kennell, David L. Soldan, and Robert W. Ives, "Using Artificial Neural Networks and Feature Saliency Techniques for Improved Iris Segmentation", *Proceedings of International Joint Conference on Neural Networks*, 2007.
26. Marcelo Mottalli, Marta Mejail and Julio Jacobo-Berlles, "Flexible Image Segmentation and Quality Assessment for Real-Time Iris Recognition", *Proceedings of International Conference on Image Processing (ICIP)*, 2009.
27. E. Levina and P.J. Bickel. Maximum likelihood estimation of intrinsic dimension. In *Advances in Neural Information Processing Systems*, volume 17, Cambridge, MA, USA, 2004. The MIT Press.
28. A.P. Papliński, "Neuro-Fuzzy Comp. — Ch. 8", pp. 8.1-8.13, May 12, 2005.
29. Dobeš, M., Martinek J., Skoupil D., Dobešová Z., Pospíšil J., Human eye localization using the modified Hough transform. *Optik*, Volume 117, No.10, p.468-473, Elsevier 2006, ISSN 0030-4026.
30. Dobeš M., Machala L., Tichavský P., Pospíšil J., Human Eye Iris Recognition Using the Mutual Information. *Optik* Volume 115, No.9, p.399-405, Elsevier 2004, ISSN 0030-4026.
31. Michal Dobeš and Libor Machala, *Iris Database*, <http://www.inf.upol.cz/iris/>
32. Kaushik Roy, Prabir Bhattacharya and Ramesh Chandra Debnath, "Multi-Class SVM Based Iris Recognition", *10th International Conference on Computer and information technology*, 2007.



GLOBAL JOURNALS INC. (US) GUIDELINES HANDBOOK 2013

WWW.GLOBALJOURNALS.ORG

FELLOW OF ASSOCIATION OF RESEARCH SOCIETY IN COMPUTING (FARSC)

- 'FARSC' title will be awarded to the person after approval of Editor-in-Chief and Editorial Board. The title 'FARSC' can be added to name in the following manner. eg. **Dr. John E. Hall, Ph.D., FARSC or William Walldroff Ph. D., M.S., FARSC**
- Being FARSC is a respectful honor. It authenticates your research activities. After becoming FARSC, you can use 'FARSC' title as you use your degree in suffix of your name. This will definitely will enhance and add up your name. You can use it on your Career Counseling Materials/CV/Resume/Visiting Card/Name Plate etc.
- 60% Discount will be provided to FARSC members for publishing research papers in Global Journals Inc., if our Editorial Board and Peer Reviewers accept the paper. For the life time, if you are author/co-author of any paper bill sent to you will automatically be discounted one by 60%
- FARSC will be given a renowned, secure, free professional email address with 100 GB of space eg.johnhall@globaljournals.org. You will be facilitated with Webmail, Spam Assassin, Email Forwarders, Auto-Responders, Email Delivery Route tracing, etc.
- FARSC member is eligible to become paid peer reviewer at Global Journals Inc. to earn up to 15% of realized author charges taken from author of respective paper. After reviewing 5 or more papers you can request to transfer the amount to your bank account or to your PayPal account.
- Eg. If we had taken 420 USD from author, we can send 63 USD to your account.
- FARSC member can apply for free approval, grading and certification of some of their Educational and Institutional Degrees from Global Journals Inc. (US) and Open Association of Research, Society U.S.A.
- After you are FARSC. You can send us scanned copy of all of your documents. We will verify, grade and certify them within a month. It will be based on your academic records, quality of research papers published by you, and 50 more criteria. This is beneficial for your job interviews as recruiting organization need not just rely on you for authenticity and your unknown qualities, you would have authentic ranks of all of your documents. Our scale is unique worldwide.
- FARSC member can proceed to get benefits of free research podcasting in Global Research Radio with their research documents, slides and online movies.
- After your publication anywhere in the world, you can upload you research paper with your recorded voice or you can use our professional RJs to record your paper their voice. We can also stream your conference videos and display your slides online.
- FARSC will be eligible for free application of Standardization of their Researches by Open Scientific Standards. Standardization is next step and level after publishing in a journal. A team of research and professional will work with you to take your research to its next level, which is worldwide open standardization.

- FARSC is eligible to earn from their researches: While publishing his paper with Global Journals Inc. (US), FARSC can decide whether he/she would like to publish his/her research in closed manner. When readers will buy that individual research paper for reading, 80% of its earning by Global Journals Inc. (US) will be transferred to FARSC member's bank account after certain threshold balance. There is no time limit for collection. FARSC member can decide its price and we can help in decision.

MEMBER OF ASSOCIATION OF RESEARCH SOCIETY IN COMPUTING (MARSC)

- 'MARSC' title will be awarded to the person after approval of Editor-in-Chief and Editorial Board. The title 'MARSC' can be added to name in the following manner. eg. Dr. John E. Hall, Ph.D., MARSC or William Walldroff Ph. D., M.S., MARSC
- Being MARSC is a respectful honor. It authenticates your research activities. After becoming MARSC, you can use 'MARSC' title as you use your degree in suffix of your name. This will definitely will enhance and add up your name. You can use it on your Career Counseling Materials/CV/Resume/Visiting Card/Name Plate etc.
- 40% Discount will be provided to MARSC members for publishing research papers in Global Journals Inc., if our Editorial Board and Peer Reviewers accept the paper. For the life time, if you are author/co-author of any paper bill sent to you will automatically be discounted one by 60%
- MARSC will be given a renowned, secure, free professional email address with 30 GB of space eg.johnhall@globaljournals.org. You will be facilitated with Webmail, Spam Assassin, Email Forwarders, Auto-Responders, Email Delivery Route tracing, etc.
- MARSC member is eligible to become paid peer reviewer at Global Journals Inc. to earn up to 10% of realized author charges taken from author of respective paper. After reviewing 5 or more papers you can request to transfer the amount to your bank account or to your PayPal account.
- MARSC member can apply for free approval, grading and certification of some of their Educational and Institutional Degrees from Global Journals Inc. (US) and Open Association of Research, Society U.S.A.
- MARSC is eligible to earn from their researches: While publishing his paper with Global Journals Inc. (US), MARSC can decide whether he/she would like to publish his/her research in closed manner. When readers will buy that individual research paper for reading, 40% of its earning by Global Journals Inc. (US) will be transferred to MARSC member's bank account after certain threshold balance. There is no time limit for collection. MARSC member can decide its price and we can help in decision.

AUXILIARY MEMBERSHIPS

ANNUAL MEMBER

- Annual Member will be authorized to receive e-Journal GJCST for one year (subscription for one year).
- The member will be allotted free 1 GB Web-space along with subDomain to contribute and participate in our activities.
- A professional email address will be allotted free 500 MB email space.

PAPER PUBLICATION

- The members can publish paper once. The paper will be sent to two-peer reviewer. The paper will be published after the acceptance of peer reviewers and Editorial Board.

PROCESS OF SUBMISSION OF RESEARCH PAPER

The Area or field of specialization may or may not be of any category as mentioned in 'Scope of Journal' menu of the GlobalJournals.org website. There are 37 Research Journal categorized with Six parental Journals GJCST, GJMR, GJRE, GJMBR, GJSFR, GJHSS. For Authors should prefer the mentioned categories. There are three widely used systems UDC, DDC and LCC. The details are available as 'Knowledge Abstract' at Home page. The major advantage of this coding is that, the research work will be exposed to and shared with all over the world as we are being abstracted and indexed worldwide.

The paper should be in proper format. The format can be downloaded from first page of 'Author Guideline' Menu. The Author is expected to follow the general rules as mentioned in this menu. The paper should be written in MS-Word Format (*.DOC, *.DOCX).

The Author can submit the paper either online or offline. The authors should prefer online submission. Online Submission: There are three ways to submit your paper:

(A) (I) First, register yourself using top right corner of Home page then Login. If you are already registered, then login using your username and password.

(II) Choose corresponding Journal.

(III) Click 'Submit Manuscript'. Fill required information and Upload the paper.

(B) If you are using Internet Explorer, then Direct Submission through Homepage is also available.

(C) If these two are not convenient, and then email the paper directly to dean@globaljournals.org.

Offline Submission: Author can send the typed form of paper by Post. However, online submission should be preferred.



PREFERRED AUTHOR GUIDELINES

MANUSCRIPT STYLE INSTRUCTION (Must be strictly followed)

Page Size: 8.27" X 11"

- Left Margin: 0.65
- Right Margin: 0.65
- Top Margin: 0.75
- Bottom Margin: 0.75
- Font type of all text should be Swis 721 Lt BT.
- Paper Title should be of Font Size 24 with one Column section.
- Author Name in Font Size of 11 with one column as of Title.
- Abstract Font size of 9 Bold, "Abstract" word in Italic Bold.
- Main Text: Font size 10 with justified two columns section
- Two Column with Equal Column with of 3.38 and Gaping of .2
- First Character must be three lines Drop capped.
- Paragraph before Spacing of 1 pt and After of 0 pt.
- Line Spacing of 1 pt
- Large Images must be in One Column
- Numbering of First Main Headings (Heading 1) must be in Roman Letters, Capital Letter, and Font Size of 10.
- Numbering of Second Main Headings (Heading 2) must be in Alphabets, Italic, and Font Size of 10.

You can use your own standard format also.

Author Guidelines:

1. General,
2. Ethical Guidelines,
3. Submission of Manuscripts,
4. Manuscript's Category,
5. Structure and Format of Manuscript,
6. After Acceptance.

1. GENERAL

Before submitting your research paper, one is advised to go through the details as mentioned in following heads. It will be beneficial, while peer reviewer justify your paper for publication.

Scope

The Global Journals Inc. (US) welcome the submission of original paper, review paper, survey article relevant to the all the streams of Philosophy and knowledge. The Global Journals Inc. (US) is parental platform for Global Journal of Computer Science and Technology, Researches in Engineering, Medical Research, Science Frontier Research, Human Social Science, Management, and Business organization. The choice of specific field can be done otherwise as following in Abstracting and Indexing Page on this Website. As the all Global

Journals Inc. (US) are being abstracted and indexed (in process) by most of the reputed organizations. Topics of only narrow interest will not be accepted unless they have wider potential or consequences.

2. ETHICAL GUIDELINES

Authors should follow the ethical guidelines as mentioned below for publication of research paper and research activities.

Papers are accepted on strict understanding that the material in whole or in part has not been, nor is being, considered for publication elsewhere. If the paper once accepted by Global Journals Inc. (US) and Editorial Board, will become the copyright of the Global Journals Inc. (US).

Authorship: The authors and coauthors should have active contribution to conception design, analysis and interpretation of findings. They should critically review the contents and drafting of the paper. All should approve the final version of the paper before submission

The Global Journals Inc. (US) follows the definition of authorship set up by the Global Academy of Research and Development. According to the Global Academy of R&D authorship, criteria must be based on:

- 1) Substantial contributions to conception and acquisition of data, analysis and interpretation of the findings.
- 2) Drafting the paper and revising it critically regarding important academic content.
- 3) Final approval of the version of the paper to be published.

All authors should have been credited according to their appropriate contribution in research activity and preparing paper. Contributors who do not match the criteria as authors may be mentioned under Acknowledgement.

Acknowledgements: Contributors to the research other than authors credited should be mentioned under acknowledgement. The specifications of the source of funding for the research if appropriate can be included. Suppliers of resources may be mentioned along with address.

Appeal of Decision: The Editorial Board's decision on publication of the paper is final and cannot be appealed elsewhere.

Permissions: It is the author's responsibility to have prior permission if all or parts of earlier published illustrations are used in this paper.

Please mention proper reference and appropriate acknowledgements wherever expected.

If all or parts of previously published illustrations are used, permission must be taken from the copyright holder concerned. It is the author's responsibility to take these in writing.

Approval for reproduction/modification of any information (including figures and tables) published elsewhere must be obtained by the authors/copyright holders before submission of the manuscript. Contributors (Authors) are responsible for any copyright fee involved.

3. SUBMISSION OF MANUSCRIPTS

Manuscripts should be uploaded via this online submission page. The online submission is most efficient method for submission of papers, as it enables rapid distribution of manuscripts and consequently speeds up the review procedure. It also enables authors to know the status of their own manuscripts by emailing us. Complete instructions for submitting a paper is available below.

Manuscript submission is a systematic procedure and little preparation is required beyond having all parts of your manuscript in a given format and a computer with an Internet connection and a Web browser. Full help and instructions are provided on-screen. As an author, you will be prompted for login and manuscript details as Field of Paper and then to upload your manuscript file(s) according to the instructions.



To avoid postal delays, all transaction is preferred by e-mail. A finished manuscript submission is confirmed by e-mail immediately and your paper enters the editorial process with no postal delays. When a conclusion is made about the publication of your paper by our Editorial Board, revisions can be submitted online with the same procedure, with an occasion to view and respond to all comments.

Complete support for both authors and co-author is provided.

4. MANUSCRIPT'S CATEGORY

Based on potential and nature, the manuscript can be categorized under the following heads:

Original research paper: Such papers are reports of high-level significant original research work.

Review papers: These are concise, significant but helpful and decisive topics for young researchers.

Research articles: These are handled with small investigation and applications.

Research letters: The letters are small and concise comments on previously published matters.

5. STRUCTURE AND FORMAT OF MANUSCRIPT

The recommended size of original research paper is less than seven thousand words, review papers fewer than seven thousands words also. Preparation of research paper or how to write research paper, are major hurdle, while writing manuscript. The research articles and research letters should be fewer than three thousand words, the structure original research paper; sometime review paper should be as follows:

Papers: These are reports of significant research (typically less than 7000 words equivalent, including tables, figures, references), and comprise:

- (a) Title should be relevant and commensurate with the theme of the paper.
- (b) A brief Summary, "Abstract" (less than 150 words) containing the major results and conclusions.
- (c) Up to ten keywords, that precisely identifies the paper's subject, purpose, and focus.
- (d) An Introduction, giving necessary background excluding subheadings; objectives must be clearly declared.
- (e) Resources and techniques with sufficient complete experimental details (wherever possible by reference) to permit repetition; sources of information must be given and numerical methods must be specified by reference, unless non-standard.
- (f) Results should be presented concisely, by well-designed tables and/or figures; the same data may not be used in both; suitable statistical data should be given. All data must be obtained with attention to numerical detail in the planning stage. As reproduced design has been recognized to be important to experiments for a considerable time, the Editor has decided that any paper that appears not to have adequate numerical treatments of the data will be returned un-refereed;
- (g) Discussion should cover the implications and consequences, not just recapitulating the results; conclusions should be summarizing.
- (h) Brief Acknowledgements.
- (i) References in the proper form.

Authors should very cautiously consider the preparation of papers to ensure that they communicate efficiently. Papers are much more likely to be accepted, if they are cautiously designed and laid out, contain few or no errors, are summarizing, and be conventional to the approach and instructions. They will in addition, be published with much less delays than those that require much technical and editorial correction.



The Editorial Board reserves the right to make literary corrections and to make suggestions to improve briefness.

It is vital, that authors take care in submitting a manuscript that is written in simple language and adheres to published guidelines.

Format

Language: The language of publication is UK English. Authors, for whom English is a second language, must have their manuscript efficiently edited by an English-speaking person before submission to make sure that, the English is of high excellence. It is preferable, that manuscripts should be professionally edited.

Standard Usage, Abbreviations, and Units: Spelling and hyphenation should be conventional to The Concise Oxford English Dictionary. Statistics and measurements should at all times be given in figures, e.g. 16 min, except for when the number begins a sentence. When the number does not refer to a unit of measurement it should be spelt in full unless, it is 160 or greater.

Abbreviations supposed to be used carefully. The abbreviated name or expression is supposed to be cited in full at first usage, followed by the conventional abbreviation in parentheses.

Metric SI units are supposed to generally be used excluding where they conflict with current practice or are confusing. For illustration, 1.4 l rather than $1.4 \times 10^{-3} \text{ m}^3$, or 4 mm somewhat than $4 \times 10^{-3} \text{ m}$. Chemical formula and solutions must identify the form used, e.g. anhydrous or hydrated, and the concentration must be in clearly defined units. Common species names should be followed by underlines at the first mention. For following use the generic name should be constricted to a single letter, if it is clear.

Structure

All manuscripts submitted to Global Journals Inc. (US), ought to include:

Title: The title page must carry an instructive title that reflects the content, a running title (less than 45 characters together with spaces), names of the authors and co-authors, and the place(s) wherever the work was carried out. The full postal address in addition with the e-mail address of related author must be given. Up to eleven keywords or very brief phrases have to be given to help data retrieval, mining and indexing.

Abstract, used in Original Papers and Reviews:

Optimizing Abstract for Search Engines

Many researchers searching for information online will use search engines such as Google, Yahoo or similar. By optimizing your paper for search engines, you will amplify the chance of someone finding it. This in turn will make it more likely to be viewed and/or cited in a further work. Global Journals Inc. (US) have compiled these guidelines to facilitate you to maximize the web-friendliness of the most public part of your paper.

Key Words

A major linchpin in research work for the writing research paper is the keyword search, which one will employ to find both library and Internet resources.

One must be persistent and creative in using keywords. An effective keyword search requires a strategy and planning a list of possible keywords and phrases to try.

Search engines for most searches, use Boolean searching, which is somewhat different from Internet searches. The Boolean search uses "operators," words (and, or, not, and near) that enable you to expand or narrow your affords. Tips for research paper while preparing research paper are very helpful guideline of research paper.

Choice of key words is first tool of tips to write research paper. Research paper writing is an art. A few tips for deciding as strategically as possible about keyword search:



- One should start brainstorming lists of possible keywords before even begin searching. Think about the most important concepts related to research work. Ask, "What words would a source have to include to be truly valuable in research paper?" Then consider synonyms for the important words.
- It may take the discovery of only one relevant paper to let steer in the right keyword direction because in most databases, the keywords under which a research paper is abstracted are listed with the paper.
- One should avoid outdated words.

Keywords are the key that opens a door to research work sources. Keyword searching is an art in which researcher's skills are bound to improve with experience and time.

Numerical Methods: Numerical methods used should be clear and, where appropriate, supported by references.

Acknowledgements: Please make these as concise as possible.

References

References follow the Harvard scheme of referencing. References in the text should cite the authors' names followed by the time of their publication, unless there are three or more authors when simply the first author's name is quoted followed by et al. unpublished work has to only be cited where necessary, and only in the text. Copies of references in press in other journals have to be supplied with submitted typescripts. It is necessary that all citations and references be carefully checked before submission, as mistakes or omissions will cause delays.

References to information on the World Wide Web can be given, but only if the information is available without charge to readers on an official site. Wikipedia and Similar websites are not allowed where anyone can change the information. Authors will be asked to make available electronic copies of the cited information for inclusion on the Global Journals Inc. (US) homepage at the judgment of the Editorial Board.

The Editorial Board and Global Journals Inc. (US) recommend that, citation of online-published papers and other material should be done via a DOI (digital object identifier). If an author cites anything, which does not have a DOI, they run the risk of the cited material not being noticeable.

The Editorial Board and Global Journals Inc. (US) recommend the use of a tool such as Reference Manager for reference management and formatting.

Tables, Figures and Figure Legends

Tables: Tables should be few in number, cautiously designed, uncrowned, and include only essential data. Each must have an Arabic number, e.g. Table 4, a self-explanatory caption and be on a separate sheet. Vertical lines should not be used.

Figures: Figures are supposed to be submitted as separate files. Always take in a citation in the text for each figure using Arabic numbers, e.g. Fig. 4. Artwork must be submitted online in electronic form by e-mailing them.

Preparation of Electronic Figures for Publication

Even though low quality images are sufficient for review purposes, print publication requires high quality images to prevent the final product being blurred or fuzzy. Submit (or e-mail) EPS (line art) or TIFF (halftone/photographs) files only. MS PowerPoint and Word Graphics are unsuitable for printed pictures. Do not use pixel-oriented software. Scans (TIFF only) should have a resolution of at least 350 dpi (halftone) or 700 to 1100 dpi (line drawings) in relation to the imitation size. Please give the data for figures in black and white or submit a Color Work Agreement Form. EPS files must be saved with fonts embedded (and with a TIFF preview, if possible).

For scanned images, the scanning resolution (at final image size) ought to be as follows to ensure good reproduction: line art: >650 dpi; halftones (including gel photographs) : >350 dpi; figures containing both halftone and line images: >650 dpi.

Color Charges: It is the rule of the Global Journals Inc. (US) for authors to pay the full cost for the reproduction of their color artwork. Hence, please note that, if there is color artwork in your manuscript when it is accepted for publication, we would require you to complete and return a color work agreement form before your paper can be published.



Figure Legends: Self-explanatory legends of all figures should be incorporated separately under the heading 'Legends to Figures'. In the full-text online edition of the journal, figure legends may possibly be truncated in abbreviated links to the full screen version. Therefore, the first 100 characters of any legend should notify the reader, about the key aspects of the figure.

6. AFTER ACCEPTANCE

Upon approval of a paper for publication, the manuscript will be forwarded to the dean, who is responsible for the publication of the Global Journals Inc. (US).

6.1 Proof Corrections

The corresponding author will receive an e-mail alert containing a link to a website or will be attached. A working e-mail address must therefore be provided for the related author.

Acrobat Reader will be required in order to read this file. This software can be downloaded

(Free of charge) from the following website:

www.adobe.com/products/acrobat/readstep2.html. This will facilitate the file to be opened, read on screen, and printed out in order for any corrections to be added. Further instructions will be sent with the proof.

Proofs must be returned to the dean at dean@globaljournals.org within three days of receipt.

As changes to proofs are costly, we inquire that you only correct typesetting errors. All illustrations are retained by the publisher. Please note that the authors are responsible for all statements made in their work, including changes made by the copy editor.

6.2 Early View of Global Journals Inc. (US) (Publication Prior to Print)

The Global Journals Inc. (US) are enclosed by our publishing's Early View service. Early View articles are complete full-text articles sent in advance of their publication. Early View articles are absolute and final. They have been completely reviewed, revised and edited for publication, and the authors' final corrections have been incorporated. Because they are in final form, no changes can be made after sending them. The nature of Early View articles means that they do not yet have volume, issue or page numbers, so Early View articles cannot be cited in the conventional way.

6.3 Author Services

Online production tracking is available for your article through Author Services. Author Services enables authors to track their article - once it has been accepted - through the production process to publication online and in print. Authors can check the status of their articles online and choose to receive automated e-mails at key stages of production. The authors will receive an e-mail with a unique link that enables them to register and have their article automatically added to the system. Please ensure that a complete e-mail address is provided when submitting the manuscript.

6.4 Author Material Archive Policy

Please note that if not specifically requested, publisher will dispose off hardcopy & electronic information submitted, after the two months of publication. If you require the return of any information submitted, please inform the Editorial Board or dean as soon as possible.

6.5 Offprint and Extra Copies

A PDF offprint of the online-published article will be provided free of charge to the related author, and may be distributed according to the Publisher's terms and conditions. Additional paper offprint may be ordered by emailing us at: editor@globaljournals.org.

You must strictly follow above Author Guidelines before submitting your paper or else we will not at all be responsible for any corrections in future in any of the way.



Before start writing a good quality Computer Science Research Paper, let us first understand what is Computer Science Research Paper? So, Computer Science Research Paper is the paper which is written by professionals or scientists who are associated to Computer Science and Information Technology, or doing research study in these areas. If you are novel to this field then you can consult about this field from your supervisor or guide.

TECHNIQUES FOR WRITING A GOOD QUALITY RESEARCH PAPER:

1. Choosing the topic: In most cases, the topic is searched by the interest of author but it can be also suggested by the guides. You can have several topics and then you can judge that in which topic or subject you are finding yourself most comfortable. This can be done by asking several questions to yourself, like Will I be able to carry our search in this area? Will I find all necessary recourses to accomplish the search? Will I be able to find all information in this field area? If the answer of these types of questions will be "Yes" then you can choose that topic. In most of the cases, you may have to conduct the surveys and have to visit several places because this field is related to Computer Science and Information Technology. Also, you may have to do a lot of work to find all rise and falls regarding the various data of that subject. Sometimes, detailed information plays a vital role, instead of short information.

2. Evaluators are human: First thing to remember that evaluators are also human being. They are not only meant for rejecting a paper. They are here to evaluate your paper. So, present your Best.

3. Think Like Evaluators: If you are in a confusion or getting demotivated that your paper will be accepted by evaluators or not, then think and try to evaluate your paper like an Evaluator. Try to understand that what an evaluator wants in your research paper and automatically you will have your answer.

4. Make blueprints of paper: The outline is the plan or framework that will help you to arrange your thoughts. It will make your paper logical. But remember that all points of your outline must be related to the topic you have chosen.

5. Ask your Guides: If you are having any difficulty in your research, then do not hesitate to share your difficulty to your guide (if you have any). They will surely help you out and resolve your doubts. If you can't clarify what exactly you require for your work then ask the supervisor to help you with the alternative. He might also provide you the list of essential readings.

6. Use of computer is recommended: As you are doing research in the field of Computer Science, then this point is quite obvious.

7. Use right software: Always use good quality software packages. If you are not capable to judge good software then you can lose quality of your paper unknowingly. There are various software programs available to help you, which you can get through Internet.

8. Use the Internet for help: An excellent start for your paper can be by using the Google. It is an excellent search engine, where you can have your doubts resolved. You may also read some answers for the frequent question how to write my research paper or find model research paper. From the internet library you can download books. If you have all required books make important reading selecting and analyzing the specified information. Then put together research paper sketch out.

9. Use and get big pictures: Always use encyclopedias, Wikipedia to get pictures so that you can go into the depth.

10. Bookmarks are useful: When you read any book or magazine, you generally use bookmarks, right! It is a good habit, which helps to not to lose your continuity. You should always use bookmarks while searching on Internet also, which will make your search easier.

11. Revise what you wrote: When you write anything, always read it, summarize it and then finalize it.



12. Make all efforts: Make all efforts to mention what you are going to write in your paper. That means always have a good start. Try to mention everything in introduction, that what is the need of a particular research paper. Polish your work by good skill of writing and always give an evaluator, what he wants.

13. Have backups: When you are going to do any important thing like making research paper, you should always have backup copies of it either in your computer or in paper. This will help you to not to lose any of your important.

14. Produce good diagrams of your own: Always try to include good charts or diagrams in your paper to improve quality. Using several and unnecessary diagrams will degrade the quality of your paper by creating "hotchpotch." So always, try to make and include those diagrams, which are made by your own to improve readability and understandability of your paper.

15. Use of direct quotes: When you do research relevant to literature, history or current affairs then use of quotes become essential but if study is relevant to science then use of quotes is not preferable.

16. Use proper verb tense: Use proper verb tenses in your paper. Use past tense, to present those events that happened. Use present tense to indicate events that are going on. Use future tense to indicate future happening events. Use of improper and wrong tenses will confuse the evaluator. Avoid the sentences that are incomplete.

17. Never use online paper: If you are getting any paper on Internet, then never use it as your research paper because it might be possible that evaluator has already seen it or maybe it is outdated version.

18. Pick a good study spot: To do your research studies always try to pick a spot, which is quiet. Every spot is not for studies. Spot that suits you choose it and proceed further.

19. Know what you know: Always try to know, what you know by making objectives. Else, you will be confused and cannot achieve your target.

20. Use good quality grammar: Always use a good quality grammar and use words that will throw positive impact on evaluator. Use of good quality grammar does not mean to use tough words, that for each word the evaluator has to go through dictionary. Do not start sentence with a conjunction. Do not fragment sentences. Eliminate one-word sentences. Ignore passive voice. Do not ever use a big word when a diminutive one would suffice. Verbs have to be in agreement with their subjects. Prepositions are not expressions to finish sentences with. It is incorrect to ever divide an infinitive. Avoid clichés like the disease. Also, always shun irritating alliteration. Use language that is simple and straight forward. put together a neat summary.

21. Arrangement of information: Each section of the main body should start with an opening sentence and there should be a changeover at the end of the section. Give only valid and powerful arguments to your topic. You may also maintain your arguments with records.

22. Never start in last minute: Always start at right time and give enough time to research work. Leaving everything to the last minute will degrade your paper and spoil your work.

23. Multitasking in research is not good: Doing several things at the same time proves bad habit in case of research activity. Research is an area, where everything has a particular time slot. Divide your research work in parts and do particular part in particular time slot.

24. Never copy others' work: Never copy others' work and give it your name because if evaluator has seen it anywhere you will be in trouble.

25. Take proper rest and food: No matter how many hours you spend for your research activity, if you are not taking care of your health then all your efforts will be in vain. For a quality research, study is must, and this can be done by taking proper rest and food.

26. Go for seminars: Attend seminars if the topic is relevant to your research area. Utilize all your resources.



27. Refresh your mind after intervals: Try to give rest to your mind by listening to soft music or by sleeping in intervals. This will also improve your memory.

28. Make colleagues: Always try to make colleagues. No matter how sharper or intelligent you are, if you make colleagues you can have several ideas, which will be helpful for your research.

29. Think technically: Always think technically. If anything happens, then search its reasons, its benefits, and demerits.

30. Think and then print: When you will go to print your paper, notice that tables are not be split, headings are not detached from their descriptions, and page sequence is maintained.

31. Adding unnecessary information: Do not add unnecessary information, like, I have used MS Excel to draw graph. Do not add irrelevant and inappropriate material. These all will create superfluous. Foreign terminology and phrases are not apropos. One should NEVER take a broad view. Analogy in script is like feathers on a snake. Not at all use a large word when a very small one would be sufficient. Use words properly, regardless of how others use them. Remove quotations. Puns are for kids, not grunt readers. Amplification is a billion times of inferior quality than sarcasm.

32. Never oversimplify everything: To add material in your research paper, never go for oversimplification. This will definitely irritate the evaluator. Be more or less specific. Also too, by no means, ever use rhythmic redundancies. Contractions aren't essential and shouldn't be there used. Comparisons are as terrible as clichés. Give up ampersands and abbreviations, and so on. Remove commas, that are, not necessary. Parenthetical words however should be together with this in commas. Understatement is all the time the complete best way to put onward earth-shaking thoughts. Give a detailed literary review.

33. Report concluded results: Use concluded results. From raw data, filter the results and then conclude your studies based on measurements and observations taken. Significant figures and appropriate number of decimal places should be used. Parenthetical remarks are prohibitive. Proofread carefully at final stage. In the end give outline to your arguments. Spot out perspectives of further study of this subject. Justify your conclusion by at the bottom of them with sufficient justifications and examples.

34. After conclusion: Once you have concluded your research, the next most important step is to present your findings. Presentation is extremely important as it is the definite medium through which your research is going to be in print to the rest of the crowd. Care should be taken to categorize your thoughts well and present them in a logical and neat manner. A good quality research paper format is essential because it serves to highlight your research paper and bring to light all necessary aspects in your research.

INFORMAL GUIDELINES OF RESEARCH PAPER WRITING

Key points to remember:

- Submit all work in its final form.
- Write your paper in the form, which is presented in the guidelines using the template.
- Please note the criterion for grading the final paper by peer-reviewers.

Final Points:

A purpose of organizing a research paper is to let people to interpret your effort selectively. The journal requires the following sections, submitted in the order listed, each section to start on a new page.

The introduction will be compiled from reference matter and will reflect the design processes or outline of basis that direct you to make study. As you will carry out the process of study, the method and process section will be constructed as like that. The result segment will show related statistics in nearly sequential order and will direct the reviewers next to the similar intellectual paths throughout the data that you took to carry out your study. The discussion section will provide understanding of the data and projections as to the implication of the results. The use of good quality references all through the paper will give the effort trustworthiness by representing an alertness of prior workings.



Writing a research paper is not an easy job no matter how trouble-free the actual research or concept. Practice, excellent preparation, and controlled record keeping are the only means to make straightforward the progression.

General style:

Specific editorial column necessities for compliance of a manuscript will always take over from directions in these general guidelines.

To make a paper clear

- Adhere to recommended page limits

Mistakes to evade

- Insertion a title at the foot of a page with the subsequent text on the next page
- Separating a table/chart or figure - impound each figure/table to a single page
- Submitting a manuscript with pages out of sequence

In every sections of your document

- Use standard writing style including articles ("a", "the," etc.)
- Keep on paying attention on the research topic of the paper
- Use paragraphs to split each significant point (excluding for the abstract)
- Align the primary line of each section
- Present your points in sound order
- Use present tense to report well accepted
- Use past tense to describe specific results
- Shun familiar wording, don't address the reviewer directly, and don't use slang, slang language, or superlatives
- Shun use of extra pictures - include only those figures essential to presenting results

Title Page:

Choose a revealing title. It should be short. It should not have non-standard acronyms or abbreviations. It should not exceed two printed lines. It should include the name(s) and address (es) of all authors.



Abstract:

The summary should be two hundred words or less. It should briefly and clearly explain the key findings reported in the manuscript-- must have precise statistics. It should not have abnormal acronyms or abbreviations. It should be logical in itself. Shun citing references at this point.

An abstract is a brief distinct paragraph summary of finished work or work in development. In a minute or less a reviewer can be taught the foundation behind the study, common approach to the problem, relevant results, and significant conclusions or new questions.

Write your summary when your paper is completed because how can you write the summary of anything which is not yet written? Wealth of terminology is very essential in abstract. Yet, use comprehensive sentences and do not let go readability for briefness. You can maintain it succinct by phrasing sentences so that they provide more than lone rationale. The author can at this moment go straight to shortening the outcome. Sum up the study, with the subsequent elements in any summary. Try to maintain the initial two items to no more than one ruling each.

- Reason of the study - theory, overall issue, purpose
- Fundamental goal
- To the point depiction of the research
- Consequences, including definite statistics - if the consequences are quantitative in nature, account quantitative data; results of any numerical analysis should be reported
- Significant conclusions or questions that track from the research(es)

Approach:

- Single section, and succinct
- As an outline of job done, it is always written in past tense
- A conceptual should situate on its own, and not submit to any other part of the paper such as a form or table
- Center on shortening results - bound background information to a verdict or two, if completely necessary
- What you account in an conceptual must be regular with what you reported in the manuscript
- Exact spelling, clearness of sentences and phrases, and appropriate reporting of quantities (proper units, important statistics) are just as significant in an abstract as they are anywhere else

Introduction:

The **Introduction** should "introduce" the manuscript. The reviewer should be presented with sufficient background information to be capable to comprehend and calculate the purpose of your study without having to submit to other works. The basis for the study should be offered. Give most important references but shun difficult to make a comprehensive appraisal of the topic. In the introduction, describe the problem visibly. If the problem is not acknowledged in a logical, reasonable way, the reviewer will have no attention in your result. Speak in common terms about techniques used to explain the problem, if needed, but do not present any particulars about the protocols here. Following approach can create a valuable beginning:

- Explain the value (significance) of the study
- Shield the model - why did you employ this particular system or method? What is its compensation? You strength remark on its appropriateness from a abstract point of vision as well as point out sensible reasons for using it.
- Present a justification. Status your particular theory (es) or aim(s), and describe the logic that led you to choose them.
- Very for a short time explain the tentative propose and how it skilled the declared objectives.

Approach:

- Use past tense except for when referring to recognized facts. After all, the manuscript will be submitted after the entire job is done.
- Sort out your thoughts; manufacture one key point with every section. If you make the four points listed above, you will need a least of four paragraphs.



- Present surroundings information only as desirable in order hold up a situation. The reviewer does not desire to read the whole thing you know about a topic.
- Shape the theory/purpose specifically - do not take a broad view.
- As always, give awareness to spelling, simplicity and correctness of sentences and phrases.

Procedures (Methods and Materials):

This part is supposed to be the easiest to carve if you have good skills. A sound written Procedures segment allows a capable scientist to replacement your results. Present precise information about your supplies. The suppliers and clarity of reagents can be helpful bits of information. Present methods in sequential order but linked methodologies can be grouped as a segment. Be concise when relating the protocols. Attempt for the least amount of information that would permit another capable scientist to spare your outcome but be cautious that vital information is integrated. The use of subheadings is suggested and ought to be synchronized with the results section. When a technique is used that has been well described in another object, mention the specific item describing a way but draw the basic principle while stating the situation. The purpose is to text all particular resources and broad procedures, so that another person may use some or all of the methods in one more study or referee the scientific value of your work. It is not to be a step by step report of the whole thing you did, nor is a methods section a set of orders.

Materials:

- Explain materials individually only if the study is so complex that it saves liberty this way.
- Embrace particular materials, and any tools or provisions that are not frequently found in laboratories.
- Do not take in frequently found.
- If use of a definite type of tools.
- Materials may be reported in a part section or else they may be recognized along with your measures.

Methods:

- Report the method (not particulars of each process that engaged the same methodology)
- Describe the method entirely
- To be succinct, present methods under headings dedicated to specific dealings or groups of measures
- Simplify - details how procedures were completed not how they were exclusively performed on a particular day.
- If well known procedures were used, account the procedure by name, possibly with reference, and that's all.

Approach:

- It is embarrassed or not possible to use vigorous voice when documenting methods with no using first person, which would focus the reviewer's interest on the researcher rather than the job. As a result when script up the methods most authors use third person passive voice.
- Use standard style in this and in every other part of the paper - avoid familiar lists, and use full sentences.

What to keep away from

- Resources and methods are not a set of information.
- Skip all descriptive information and surroundings - save it for the argument.
- Leave out information that is immaterial to a third party.

Results:

The principle of a results segment is to present and demonstrate your conclusion. Create this part a entirely objective details of the outcome, and save all understanding for the discussion.

The page length of this segment is set by the sum and types of data to be reported. Carry on to be to the point, by means of statistics and tables, if suitable, to present consequences most efficiently. You must obviously differentiate material that would usually be incorporated in a study editorial from any unprocessed data or additional appendix matter that would not be available. In fact, such matter should not be submitted at all except requested by the instructor.



Content

- Sum up your conclusion in text and demonstrate them, if suitable, with figures and tables.
- In manuscript, explain each of your consequences, point the reader to remarks that are most appropriate.
- Present a background, such as by describing the question that was addressed by creation an exacting study.
- Explain results of control experiments and comprise remarks that are not accessible in a prescribed figure or table, if appropriate.
- Examine your data, then prepare the analyzed (transformed) data in the form of a figure (graph), table, or in manuscript form.

What to stay away from

- Do not discuss or infer your outcome, report surroundings information, or try to explain anything.
- Not at all, take in raw data or intermediate calculations in a research manuscript.
- Do not present the similar data more than once.
- Manuscript should complement any figures or tables, not duplicate the identical information.
- Never confuse figures with tables - there is a difference.

Approach

- As forever, use past tense when you submit to your results, and put the whole thing in a reasonable order.
- Put figures and tables, appropriately numbered, in order at the end of the report
- If you desire, you may place your figures and tables properly within the text of your results part.

Figures and tables

- If you put figures and tables at the end of the details, make certain that they are visibly distinguished from any attach appendix materials, such as raw facts
- Despite of position, each figure must be numbered one after the other and complete with subtitle
- In spite of position, each table must be titled, numbered one after the other and complete with heading
- All figure and table must be adequately complete that it could situate on its own, divide from text

Discussion:

The Discussion is expected the trickiest segment to write and describe. A lot of papers submitted for journal are discarded based on problems with the Discussion. There is no head of state for how long a argument should be. Position your understanding of the outcome visibly to lead the reviewer through your conclusions, and then finish the paper with a summing up of the implication of the study. The purpose here is to offer an understanding of your results and hold up for all of your conclusions, using facts from your research and generally accepted information, if suitable. The implication of result should be visibly described. Infer your data in the conversation in suitable depth. This means that when you clarify an observable fact you must explain mechanisms that may account for the observation. If your results vary from your prospect, make clear why that may have happened. If your results agree, then explain the theory that the proof supported. It is never suitable to just state that the data approved with prospect, and let it drop at that.

- Make a decision if each premise is supported, discarded, or if you cannot make a conclusion with assurance. Do not just dismiss a study or part of a study as "uncertain."
- Research papers are not acknowledged if the work is imperfect. Draw what conclusions you can based upon the results that you have, and take care of the study as a finished work
- You may propose future guidelines, such as how the experiment might be personalized to accomplish a new idea.
- Give details all of your remarks as much as possible, focus on mechanisms.
- Make a decision if the tentative design sufficiently addressed the theory, and whether or not it was correctly restricted.
- Try to present substitute explanations if sensible alternatives be present.
- One research will not counter an overall question, so maintain the large picture in mind, where do you go next? The best studies unlock new avenues of study. What questions remain?
- Recommendations for detailed papers will offer supplementary suggestions.

Approach:

- When you refer to information, differentiate data generated by your own studies from available information
- Submit to work done by specific persons (including you) in past tense.
- Submit to generally acknowledged facts and main beliefs in present tense.



ADMINISTRATION RULES LISTED BEFORE SUBMITTING YOUR RESEARCH PAPER TO GLOBAL JOURNALS INC. (US)

Please carefully note down following rules and regulation before submitting your Research Paper to Global Journals Inc. (US):

Segment Draft and Final Research Paper: You have to strictly follow the template of research paper. If it is not done your paper may get rejected.

- The **major constraint** is that you must independently make all content, tables, graphs, and facts that are offered in the paper. You must write each part of the paper wholly on your own. The Peer-reviewers need to identify your own perceptive of the concepts in your own terms. NEVER extract straight from any foundation, and never rephrase someone else's analysis.
- Do not give permission to anyone else to "PROOFREAD" your manuscript.
- **Methods to avoid Plagiarism is applied by us on every paper, if found guilty, you will be blacklisted by all of our collaborated research groups, your institution will be informed for this and strict legal actions will be taken immediately.)**
- To guard yourself and others from possible illegal use please do not permit anyone right to use to your paper and files.



CRITERION FOR GRADING A RESEARCH PAPER (COMPILATION)
BY GLOBAL JOURNALS INC. (US)

Please note that following table is only a Grading of "Paper Compilation" and not on "Performed/Stated Research" whose grading solely depends on Individual Assigned Peer Reviewer and Editorial Board Member. These can be available only on request and after decision of Paper. This report will be the property of Global Journals Inc. (US).

Topics	Grades		
	A-B	C-D	E-F
Abstract	Clear and concise with appropriate content, Correct format. 200 words or below	Unclear summary and no specific data, Incorrect form Above 200 words	No specific data with ambiguous information Above 250 words
Introduction	Containing all background details with clear goal and appropriate details, flow specification, no grammar and spelling mistake, well organized sentence and paragraph, reference cited	Unclear and confusing data, appropriate format, grammar and spelling errors with unorganized matter	Out of place depth and content, hazy format
Methods and Procedures	Clear and to the point with well arranged paragraph, precision and accuracy of facts and figures, well organized subheads	Difficult to comprehend with embarrassed text, too much explanation but completed	Incorrect and unorganized structure with hazy meaning
Result	Well organized, Clear and specific, Correct units with precision, correct data, well structuring of paragraph, no grammar and spelling mistake	Complete and embarrassed text, difficult to comprehend	Irregular format with wrong facts and figures
Discussion	Well organized, meaningful specification, sound conclusion, logical and concise explanation, highly structured paragraph reference cited	Wordy, unclear conclusion, spurious	Conclusion is not cited, unorganized, difficult to comprehend
References	Complete and correct format, well organized	Beside the point, Incomplete	Wrong format and structuring



INDEX

A

Abundance · 2
Accuracy · 35, 46, 47, 48
Accurate · 30
Algorithm · 1, 3, 4, 5, 6, 20, 23, 26, 30, 44
Annotations · 15
Approximated · 7

B

Binary · 5, 36, 37
Biometric · 34, 36
Boundary · 33, 37, 39, 41
Brained · 26, 28
Brightness · 13, 20, 26, 39

C

Captured · 5
Clustered · 28
Collating · 13
Conference · 24, 30, 48, 49, 50
Corrupted · 1, 2, 3, 4, 8, 11

D

Deblurring · 3
Deficiency · 15, 18
Detector · 4, 5, 16, 20, 21, 24
Dimension · 37, 39, 41, 43, 45, 46, 48
Dimensional · 16, 21, 26, 33, 35, 41, 43, 44, 45, 48

E

Emphasizes · 3
Evidence · 22
Experimental · 1, 2, 3, 4, 13, 24, 45
Experimental · 3, 4, 5, 47
Exponential · 44
Extraction · 22, 23, 35
Extractors · 15

F

Filtering · 1, 3, 4, 5, 9, 13, 23

G

Goal · 1, 43
Gradients · 27
Graphs · 8, 28

H

Hexagonal · 45
Homogeneity · 39

I

Illumination · 16, 20, 27
Implemented · 28, 30
Implicit · 15
Information · 1, 3, 5, 6, 13, 15, 16, 18, 20, 22, 23, 24, 26, 28, 33, 38, 39, 40, 41, 47, 50
Irrelevant · 13, 46

M

Methodologies · 7, 34
Multidomain · 35, 48
Multiple · 26, 35

O

Occurrence · 39
Optimization · 35, 36, 37, 43, 47

P

Pathology · 16
Preservation · 4, 9
Probability · 39

Q

Quantitatively · 2

R

Recognition · 20, 33, 34, 35, 36, 41, 48, 49, 50

Recognition · 48, 49, 50

Relevance · 18, 23

Resonance · 1, 18

S

Satellite · 1

Shape · 35, 36, 37, 39, 49

Simplest · 7, 21

Sparse · 46, 47

Springer · 9

Symmetric · 40, 41

T

Technique · 3, 5, 13, 18, 26, 33, 34, 35, 36, 37, 41, 45, 46

Transmission · 1

U

Underneath · 33, 35

V

Visually · 2, 13, 23

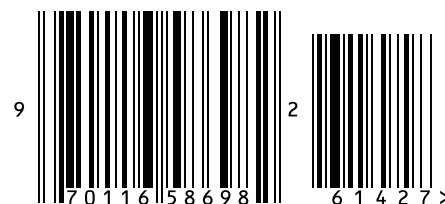


save our planet



Global Journal of Computer Science and Technology

Visit us on the Web at www.GlobalJournals.org | www.ComputerResearch.org
or email us at helpdesk@globaljournals.org



ISSN 9754350