

GLOBAL JOURNAL

OF COMPUTER SCIENCE AND TECHNOLOGY: F

Graphics & Vision

CMOS Image Sensors

Robust Audio Watermarking

Highlights

Performance Analysis

Human Skin Detection

Discovering Thoughts, Inventing Future

VOLUME 13

ISSUE 3

VERSION 10



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

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

VOLUME 13 ISSUE 3 (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: Open Association of Research Society
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. Image Toolbox for CMOS Image Sensors Fast Simulation. ***1-6***
 - 2. Digital Image Watermarking using DCT and ZIP Compression Technique. ***7-10***
 - 3. Performance Analysis among Different Classifier Including Naive Bayes, Support Vector Machine and C4.5 for Automatic Weeds Classification. ***11-15***
 - 4. Software for Drawing Jittered Plot: A Graphical Counterpart of Anova. ***17-24***
 - 5. Data Adaptive Multi-Band Approach for Robust Audio Watermarking using Adaptive Threshold. ***25-30***
 - 6. Human Skin Detection. ***31-37***
-
- vii. Auxiliary Memberships
 - viii. Process of Submission of Research Paper
 - ix. Preferred Author Guidelines
 - x. Index



Image Toolbox for CMOS Image Sensors Fast Simulation

By David Navarro, Zhenfu Feng & Ian O'Connor

Université de Lyon

Abstract - This paper presents a new toolbox, dedicated to image sensors designers. It permits fast analog simulations with a novel parametric simulation method. Image sensor matrixes are composed of millions of pixels. Such architectures are long to simulate, if not impossible; and input/ output data management is complex. A graphical toolbox has been developed in Cadence Analog Design Environment (ADE) to overcome these problems. That toolbox has been completely written in Cadence SKILL language. It permits to manipulate images as input, launch parametric analysis on a single pixel, and interpolate results for all pixels in the matrix. It then automatically generates images at output, in order to check the quality of microelectronic design. Low level aspects can also be analyzed at system (image) level.

Keywords : *image sensor, microelectronic design, APS, simulation, modeling.*

GJCST-F Classification: *1.4.8*



IMAGE TOOLBOX FOR CMOS IMAGE SENSORS FAST SIMULATION

Strictly as per the compliance and regulations of:



Image Toolbox for CMOS Image Sensors Fast Simulation

David Navarro ^α, Zhenfu Feng ^σ & Ian O'Connor ^ρ

Abstract - This paper presents a new toolbox, dedicated to image sensors designers. It permits fast analog simulations with a novel parametric simulation method. Image sensor matrixes are composed of millions of pixels. Such architectures are long to simulate, if not impossible; and input/output data management is complex. A graphical toolbox has been developed in Cadence Analog Design Environment (ADE) to overcome these problems. That toolbox has been completely written in Cadence SKILL language. It permits to manipulate images as input, launch parametric analysis on a single pixel, and interpolate results for all pixels in the matrix. It then automatically generates images at output, in order to check the quality of microelectronic design. Low level aspects can also be analyzed at system (image) level.

Keywords : image sensor, microelectronic design, APS, simulation, modeling.

I. INTRODUCTION

Image sensors in standard CMOS technology are now a well established alternative to the CCD image sensors technology. Indeed, process maturation, integration possibilities and low power consumption make CMOS image sensor widespread. Moreover, recent 3D technologies focused researchers and industry on new image sensor architectures and 3D floorplanning [1] [2].

CMOS image sensors are electronic systems that are composed of analog blocks (pixel matrix, noise reduction block (CDS: Correlated Double Sampling), column amplifier), digital blocks (controller and decoders), and mixed blocks (ADC: Analog to Digital Converter) [3]. Fig. 1 details a basic CMOS image sensor floorplan.

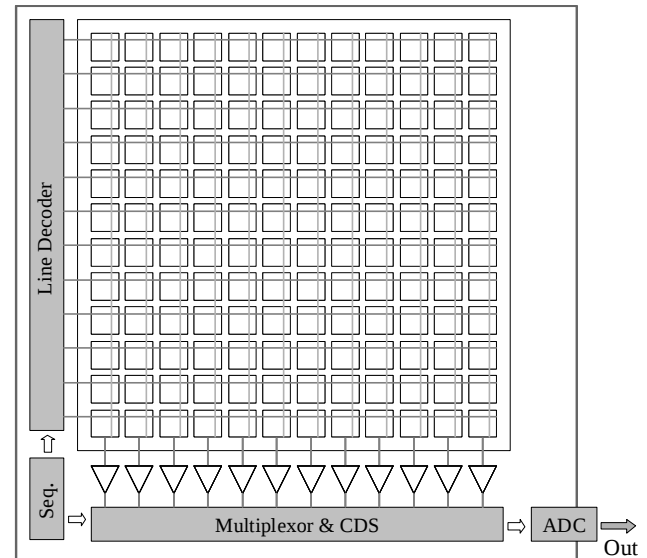


Figure 1 : Classical CMOS image sensor floorplan

II. PROBLEMS IN IMAGE SENSOR SIMULATIONS AND STATE OF ART

A first major problem in these matrix structures is density: it is difficult to make system-level analysis because of the multiplicity of inputs and outputs. Indeed, in a classical pixel, a 3-transistors active pixel, also called 3-T APS [3], shown in Fig.2, three inputs are necessary. A "reset" signal –for initialization-, a "select" signal –for reading- and luminosity are required. The two digital control signals, "reset" and "select", have precise timings. These signals are common for each line, and have to respect an interlaced sequence. In a M x N matrix, M lines have to be considered to generate these signals. The major problem comes in fact from the third input: M x N luminosity inputs have to be considered.

Author ^α : Université de Lyon; Institut des Nanotechnologies de Lyon INL-UMR5270, CNRS, Ecole Centrale de Lyon, Ecully, F-69134, France. E-mail : david.navarrohere@eclyon.fr

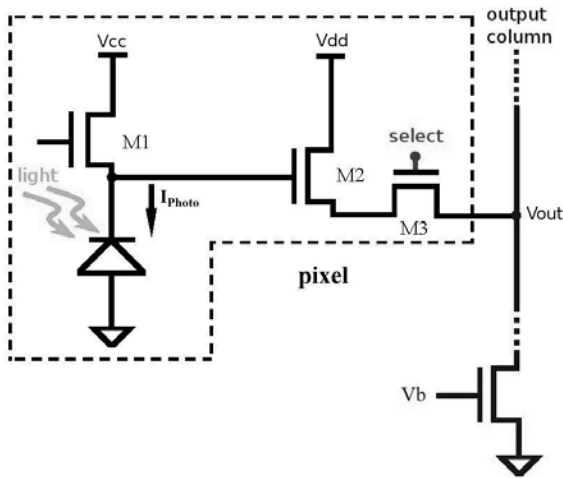


Figure 2 : Classical CMOS 3-T Active Pixel Sensor (APS)

Luminosity is often set as a constant value, anyway classical pixels integrate current over an exposure time (also called integration time). Depending on the photodiode model, the luminosity is set with a current source that emulates the photocurrent if a simple equivalent schematic is drawn, or luminosity can be considered if model is of upper level (for example in Verilog-A or VHDL-AMS language [4]). Of course, it is mandatory to simulate different values on pixel to realistically simulate the sensor characteristics. The problem is how to set easily thousands or millions of design variables. The same problem exists at sensor output: it is difficult to manage and analyze millions of analog values. In classical CMOS image sensors, outputs are voltages, sometimes ADC output. The second major problem is simulation time of image sensors, as Fig. 3 shows.

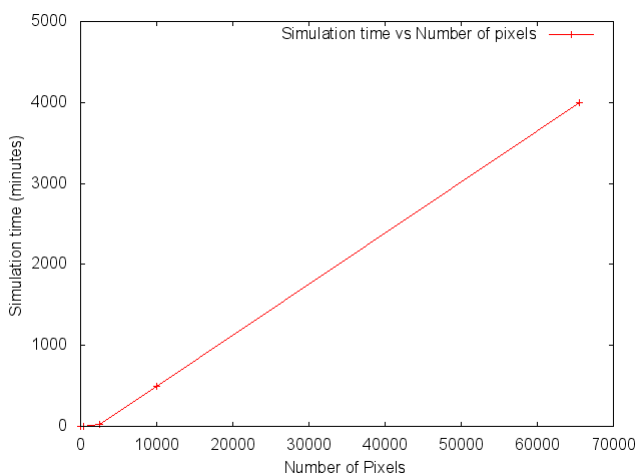


Figure 3 : Simulation time of images sensors with SPICE classical simulations

As a consequence, designers rarely simulate full imagers at a time. Blocks are validated separately, and small matrixes are used to validate the behavior.

Testchips (ASICs) are also required to properly characterize pixels characteristics. This toolbox is proposed to help designers running "image sensor – level" simulations. Available image sensor simulators are TCAD and focus on physical and FDTD simulations [5] [6]. On the other hand, image processing toolbox exist [7]. ECAD image sensor simulators are also missing.

Some high-level models –for example VHDL-AMS or MATLAB- have been developed [5], but it appears that the gap between a real analog structure and a high-level model make it uneasy to use within an optimization work. Experience showed that analog designers trust better in a level 53 SPICE model than in a third-party high-level model. If modeling and simulation tools are not integrated in classical framework, they cause another drawback: the use of several design and simulation platforms. Computer power calculation has rapidly increased past ten years, and computer clusterization is now a widespread solution, so simulation time can be lowered. Meanwhile, this paper shows that only small matrixes can be simulated with a classical approach. Moreover, designers cannot do without post-layout simulations that are easily accessible in the classical micro-electronic design flow.

To answer the above mentioned problems, we propose in this paper a graphical toolbox that is dedicated to electronic designers, as it is integrated in Cadence Analog Design Environment (ADE). Novelty is to propose an easy way to manage many (millions) parameters at input and output, and to fasten simulation with a new parametric approach.

III. INPUT AND OUTPUT IMAGE PROCESS

In order to handle all the input and output signals, a high-level mechanism has been set. It permits to read an image as input, and generates an image as output. As Fig. 4 shows, several steps are required.

We consider an image that has the same resolution as the image sensor. In that way, a pixel in the input image will be converted into luminosity or a photocurrent value. Then, that input value is set to the pixel input. A first skill processing function converts image values into lux, watts or amperes (according to the photodiode model), and then a second one assigns each value to each pixel. Design variables in ADE are used to make that mapping. Before assignment, an automatic creation of millions of design variables is done. Names are also fixed by software, for example lum_0_0 for the first pixel, lum_3_200 for 201th pixel of fourth line. Then, the simulation that was configured in ADE is run.

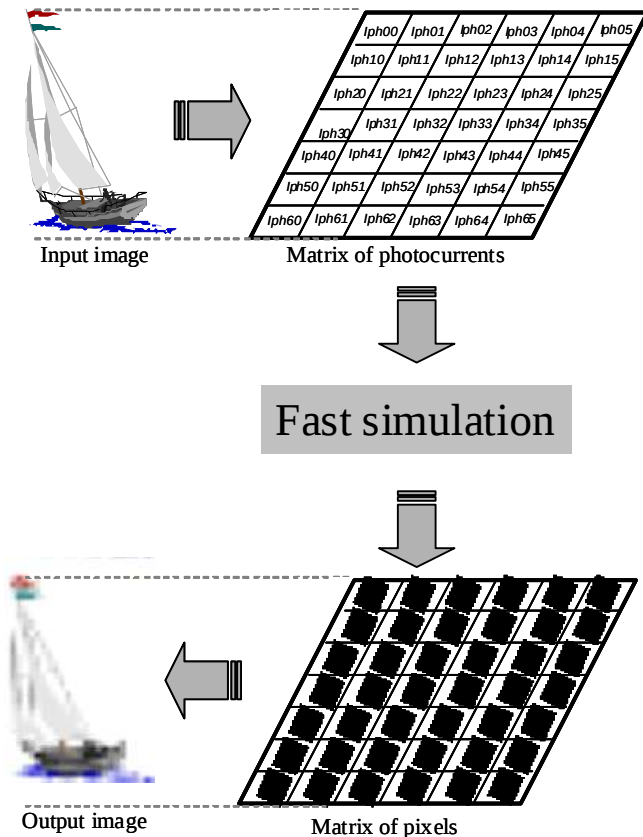


Figure 4 : Input and output image flow

As a result, pixel output voltage can be a raw voltage value or a simulated double sampled one, as illustrated in Fig. 6. In case of automatic double measurement, two simulations are run. The original netlist is modified in order to turn on the select transistor, so V_{rst} can be read at reset state. It is effectively sampled $1\mu s$ before falling edge of reset signal. Considering classical timings in such circuits, pixel output voltage will have a stable and final value. This timing is a parameter that can be changed through a setup menu. Then, the user-defined simulation is run, and V_{os} signal is sampled at the middle time of select window.

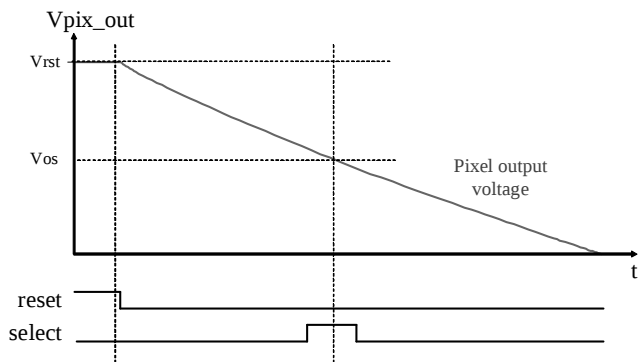


Figure 5 : Automatic timing to read voltages values at pixel output: on user request, simulator can output V_{os} or $(V_{rst} - V_{os})$ based data

At end of simulation, Cadence Spectre output files are processed in order to display results. It is mandatory to detail a classical output simulation to best explain the internal result processing. The output voltage, computed with one of the two previously explained manners, is converted in grey level. Three solutions can also be configured: raw, absolute or relative coding.

- Raw coding is a basic reading of V_{os} output signal voltage and supply voltage V_{cc} is set as maximal digital code (255 in 8-bit or 1023 in 10-bit). Equation 1 gives a calculation example for 8-bit resolution. This solution can be used as a debug mode for designer, in order to check the raw simulation output compared to classical (manual) simulation.

$$\text{Raw grey-value} = 255 \times \left(1 - \frac{V_{os}}{V_{cc}} \right) \quad (1)$$

- Absolute coding consists in setting the maximal pixel output (V_{rst}) as maximal digital code. Saturation, ie 0v, is set as the minimal value. Equation 2 gives the equivalent calculation for 8-bit resolution. This calculation gives an equivalent result as a classical CDS block would do.

$$\text{Absolute coding grey-value} = 255 \times \left(1 - \frac{V_{os}}{V_{rst}} \right) \quad (2)$$

- Relative coding consists in finding maximal and minimal pixel output voltages (respectively V_{max} and V_{min}) in the matrix, and to set them to maximal and minimal codes. Equation 3 gives the equivalent calculation for 8-bit resolution. It is what an image signal processor would do to optimize dynamic range.

$$\text{Relative coding grey-value} = 255 \times \left(1 - \frac{V_{os}}{V_{max} - V_{min}} \right) \quad (3)$$

These three solutions permit to optimize hardware and software in image sensor: hardware blocks such as CDS and ADC, and software for image signal processor at output. Hardware block that can be hierarchically optimized are at pixel, matrix, matrix and CDS, and matrix and CDS and ADC levels.

IV. FAST SIMULATION METHOD AND USER INTERFACE

Classical simulations are long, and permit only small matrixes analysis. It is to notice that very complex equations will be simulated to calculate voltages that will enter an analog to digital converter. Output will also have a discrete step, as output is composed of 256 grey levels. Instead of downsizing result at output (infinite values of voltages into 256 finite grey level values), we discretize input.

We also run a set of 256 equations that composes all the possible inputs on a single pixel, whose line and column capacitances have been considered to render the matrix real values. All the possible outputs are also known. This permits to create a lookup table, as shown in Fig. 6.

255	↔	1.05v	↔	121
254	↔	1.06v	↔	120
253	↔	1.07v	↔	119
252	↔	1.08v	↔	117
⋮				
2	↔	1.81v	↔	24
1	↔	1.82v	↔	23
0	↔	1.83v	↔	21

Look Up Table

Figure 6 : Look up table that associates input grey levels to output pixel voltage and then output grey levels

After this 256 simulations step (for look up table establishment), each of the MxN inputs are associated to an output during a mapping step.

The graphical user interface is showed in Fig. 7. It is opened via an "imager" menu in ADE.

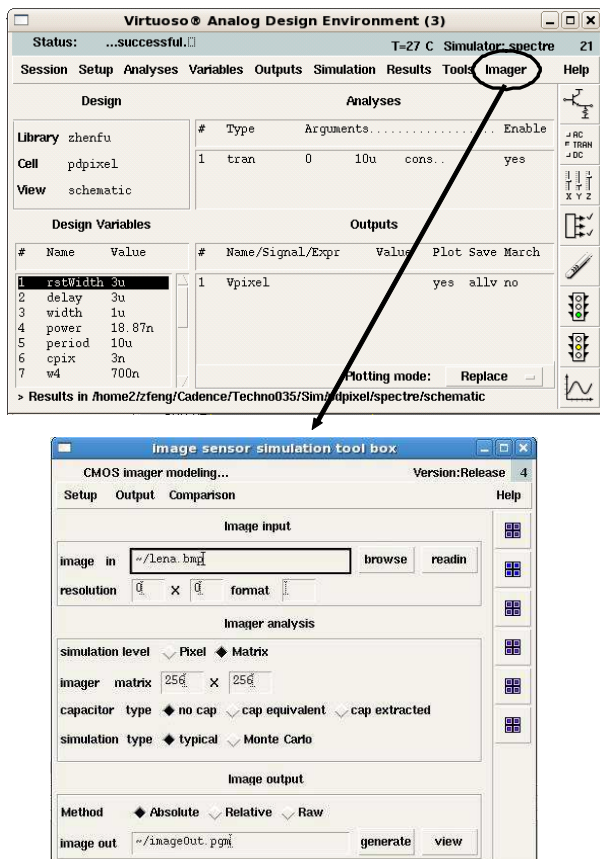


Figure 7 : Fast simulator user interface

User can select the input image (in uncompressed BMP or PGM formats). Clicking the read-in button imports image and converts color or grey pixels values into light information (light power or photocurrent according to the photodiode model that is used). This conversion is done by considering silicon sensitivity to light.

As analysis setup is launched from a single pixel schematic, it is possible to select a single pixel simulation (sim-Pixel), or a matrix simulation (sim-Matrix). For a single pixel simulation, coordinates of pixel within the matrix (pixel at location 100, 100 in Fig. 7) are selected. For a m x n matrix simulation, the same pixel is simulated m x n times with respective m x n stimuli. As a matter of fact, output node of the pixel (vout) as to be defined. Vout is set as "pixel" value in Fig. 7. It is to notice that many pixel structures can be simulated, since simulation and toolbox can run on any schematic. Moreover, this toolbox could be used for any matrix simulation, like memories.

V. RESULTS

As test example, we have considered an input image, a classical 3T pixel, and several configurations. As we can observe, from an input image, Fig. 8, it is possible to automatically set the matrix data (light) input that is applied on electronic structure, to simulate electronic circuit with Cadence Spectre, and then to monitor output images according to the chosen reading mode configuration (Fig. 9 to Fig. 11). Raw or absolute coding will be used to characterize low level characteristics of pixels matrix, or to study a system composed of matrix and Correlated Double Sampling block, eventually with Analog to Digital Converter (ADC).



Figure 8 : Input image and histogram

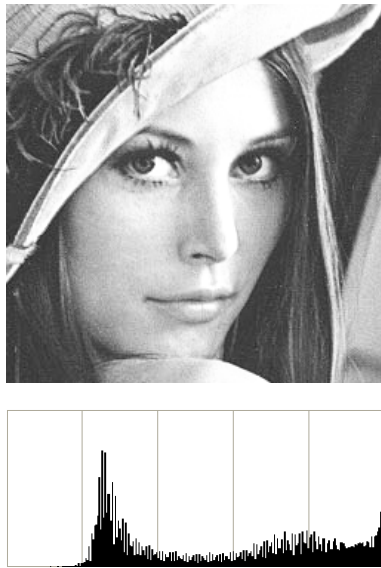


Figure 9 : Absolute coding output image and histogram,
Wfollower = 1 μm

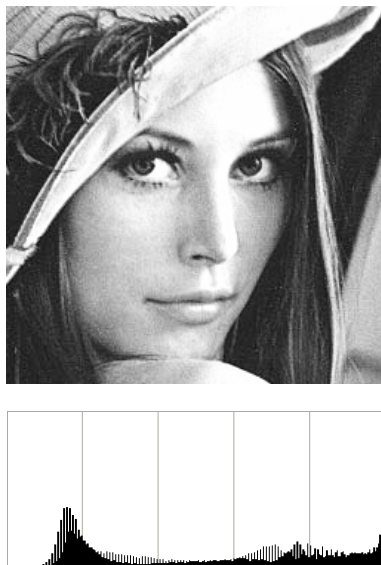


Figure 10 : Relative coding output image and histogram,
Wfollower = 1 μm

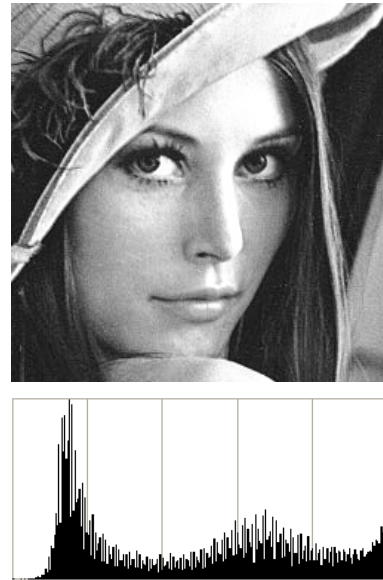


Figure 11 : Relative coding output image and histogram,
Wfollower = 15 μm

As double sampling is done in CDS block, it is indeed useless to simulate an image sensor that comprises such a block with the relative coding option. Relative coding is useful to characterize a standalone pixels matrix, or specific smart image sensors [8] that don't embed CDS. Following blocks parameters, such as amplifiers, CDS and ADC, can also be studied and dimensioned. We can observe that contrast and mean values differ from input to output images, and it is more visible on histograms. Histograms, that present grey values versus number of pixels, clearly show that transfer function of sensor is not linear. Moreover, Fig. 11 shows the transistor size impact on output image quality. In this testcase, a bigger W size gives a closer histogram so a tradeoff sensor size / quality can be led. Histogram values of Fig. 11 are shifted to the right compared to Fig. 10, this sensor suffers also a nonlinearity. Moreover, it has saturation (too light values at right of histogram). Electronic and architecture impacts on image sensor quality can also be clearly studied.

Simulation time is detailed in Table 1, gain is significant. Classical simulation is not possible on an Intel Xeon processor 2.6GHz with 4GB of RAM for a 512x512 matrix, because of memory overflow. It is to note that fast simulation time is almost constant, because look up table establishment (i.e 256 simulations) is important and look up table reading (mapping) is negligible.

Matrix size	8x8	12x12	60x60	256X256	512X512
Classical simulation (min)	3.7	7.5	179.23	3465	error
Fast simulation (min)	11.57	12.01	12.21	12.3	12.31

Table 1 : Simulation time for classical simulation and fast simulation

In terms of accuracy, table 2 gives difference between a 20x20 real matrix simulation and the one-pixel-based fast simulation.

	Pixel voltage output	error
Classical simulation	779,92527 mV	0,00043 %
Fast simulation	779,92185 mV	

Table 2 : Accuracy of fast simulation method

VI. CONCLUSION

This paper presents a new toolbox for CMOS image sensor simulations in Cadence Analog Design Environment. It is written in SKILL language, and it permits to input automatically an input image, fitting each image pixel on each pixel of the electronic sensor under study. Light stimuli is calculated and applied on each. Output images permit to check the quality of analog design in the pixel. Low level aspects of the sensor can be analyzed at system (image) level. This toolbox is composed of a new fast simulator that permits several mega-pixels simulation within a few minutes. Its principle is to study a single pixel in a parametrical way, and to project result on the matrix.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Knickerbocker, J.U. et al. "3-D Silicon Integration and Silicon Packaging Technology Using Silicon Through-Vias", IEEE Journal of Solid-State Circuits, vol. 41, no. 8, pp.1718-1725, August 2006.
2. Topol et. al., "Three-dimensional integrated circuits", IBM Journal Research and Development,. Vol. 50 no. 4/5, pp. 491-506, July/September 2006.
3. E.R. Fossum, "CMOS Image Sensors: Electronic Camera-On-A-Chip", IEEE Trans. on Electronic Devices, Vol 44, N° 10, October 1997.
4. D. Navarro, D. Ramat, F. Mieyeville, I. O'Connor, F. Gaffiot, L. Carrel, "VHDL & VHDL-AMS modeling and simulation of a CMOS imager IP", Forum on specification & Design Languages, Lausanne, Switzerland, September 2005.
5. Silvaco, "CMOS Image Sensor Simulation - ATLAS", www.silvaco.com/content/kbase/CIS_april2010.pdf
6. CrossLight Inc, "3D Simulation of CMOS Image Sensor", www.crosslight.com/applications/crosslight_3DCIS3.pdf

7. Matworks, "MATLAB Image Processing Toolbox", <http://www.mathworks.com/products/image>
8. J. Kramer, R. Sarpeshkar, C. Koch, "Pulse-Based Analog Velocity Sensors", IEEE Trans. On Circuits and Systems II : Analog and Digital Signal Processing, Vol 44, pp 86-101, 1997.



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
GRAPHICS & VISION

Volume 13 Issue 3 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

Digital Image Watermarking using DCT and ZIP Compression Technique

By Raunak Phade, Rucha Mahajan, Sampurna Mankar & Pratik Patil

Singhad College of Engineering, Pune, India

Abstract - With the exponential growth of internet and high speed networks operating throughout the world it is a challenging task to protect copyright of an individual's creation. In this paper, a spread-spectrum-like discrete cosine transform (DCT) domain watermarking technique for copyright protection of still digital images is analyzed. In our approach, the watermark is embedded in the mid frequency band of the DCT blocks only in the sub band which is carrying low frequency components and the high frequency sub band components remain untouched. With this approach we are developing an application for copyright protection and X-ray medical application too. For further security we are using ZIP Compression Technique. In effect, the main purpose of image compression is to trim down redundancy of the image data in order to be able to store or transmit data in an efficient form.

Keywords : DCT, compression, watermark, huffman.

GJCST-F Classification: I.4.2



Strictly as per the compliance and regulations of:



Digital Image Watermarking using DCT and ZIP Compression Technique

Raunak Phade^α, Rucha Mahajan^σ, Sampurna Mankar^ρ & Pratik Patil^ω

Abstract - With the exponential growth of internet and high speed networks operating throughout the world it is a challenging task to protect copyright of an individual's creation. In this paper, a spread-spectrum-like discrete cosine transform (DCT) domain watermarking technique for copyright protection of still digital images is analyzed. In our approach, the watermark is embedded in the mid frequency band of the DCT blocks only in the sub band which is carrying low frequency components and the high frequency sub band components remain untouched. With this approach we are developing an application for copyright protection and X-ray medical application too. For further security we are using ZIP Compression Technique. In effect, the main purpose of image compression is to trim down redundancy of the image data in order to be able to store or transmit data in an efficient form.

Keywords : DCT, compression, watermark, huffman.

I. INTRODUCTION

The recent development of network multimedia systems & explosion of fast communication network has discouraged & damped the multimedia content providers i.e. authors, publishers to grant the distribution of their document on the network environment. The intellectual property authentication has become an issue of concern. Also a considerable efforts and a lot of economic resources are used for the creation of intellectual property particularly in industrial society. The cost of reproducing such intellectual creations typically consist only a small portion of the creation. However a creator always desires some rewards or incentives for his creation which he is not able to get due to low cost copying. More and more researchers are attracted to the area of image watermarking because of the property of the image as it has a lot of redundant information contained in it which can be exploited to be used for watermark insertion. In the same way, image compression is very important for efficient transmission and storage of images. Image compression standards bring about many advantages, such as: (1) easy exchange of image files between different devices and applications; (2) reuse of the existing hardware and software; (3) existence of benchmarks and reference data sets, etc.

II. DISCRETE COSINE TRANSFORM (DCT)

The DCT process is applied on blocks of 8×8 which will convert into series of coefficients which define spectral composition of the block. DCT allows an image to be broken into different frequency bands i.e. the high, middle and low frequency bands as shown in Fig. 2.1 thus making it easy to choose the band in which watermark is to be inserted. The transformer also transforms the input data into a format to reduce interpixel redundancies in the input image. Transform coding techniques use linear mathematical transform to map the pixel values onto a set of coefficients. The main reason behind the success of transform-based coding schemes is that many of the resulting coefficients for most natural images have small magnitudes and can be quantized without causing significant distortion in decoded image. Due to all these reasons, the Discrete Cosine Transform (DCT) has become the most widely used transform coding techniques.

a) Watermark Embedding Algorithm

- i. Segment the image $I(i, j)$ into two sub band blocks which is half the size of the original image i.e. $I1(i/2, j)$ & $I2(i/2, j)$. Here, $I1(i/2, j)$ gives the high intensity pixels block and $I2(i/2, j)$ give low intensity pixels block $I(i, j) = \Sigma I(i/2, j) + I(i/2, j)$.
- ii. Then break the $I1(i, j/2)$ into blocks of size 8×8 .
- iii. Find the DCT (Discrete Cosine Transform) of each of the block.
- iv. Private Key is used to generate pseudo random number sequences of domain $\{-1, 0, 1\}$.
- v. Preprocess the watermark by converting the watermark into a binary sequence i.e. $W(m \times n) \rightarrow W(s \times 1)$ where $s = m \times n$.
- vi. Embed that watermark on each of the DCT block in the mid band of each coefficient block by using the pseudo random number sequence along with the watermark sequence.
- vii. After embedding the watermark the inverse DCT operation is done on the sub band so as to get the averaged image band again.

b) Watermark Extraction Algorithm

- i. Repeat step I and II of watermark embedding algorithm.
- ii. Using the private key extract watermark by applying (IDCT) Inverse Discrete Cosine transform.

Author ^{α σ ρ ω} : Department of Computer Engineering, Singhad College of Engineering, Pune, India.
E-mails : raunak_phade@rediffmail.com, m.rucha2@yahoo.com, sampurnamankar@yahoo.com, pratik1789@gmail.com

- iii. Then compare the watermark with original watermark.
- iv. Similar watermark will prove the authenticity.

III. COMPRESSION

Image compression is reducing the size in bytes of graphics file without debasing the quality of image to an unacceptable level. Also different quantization and coding techniques are used for different sub bands based on their statistical properties. By using their proposed scheme one can accomplish satisfactory reconstructed images with large compression ratios.

a) Some of the basic Compression Techniques

i. The JPEG Compression

JPEG i.e. Joint Photographic Expert Group compression can be used in a variety of file formats like as given below:

- EPS-files
- EPS DCS-files
- JFIF-files
- PDF-files

The image is partitioned into non-overlapping 8*8 blocks. After that DCT (Discrete Cosine transform) is applied to each block so as to convert the spatial domain gray level of pixels into coefficients in frequency domain.

After computation of DCT coefficients is done, they are normalized according to a quantization table with different scales provided by the JPEG standard computed. Then the Quantized coefficients are rearranged into a zigzag scan order as shown in Fig. 2.2. The process may be acquired as-

1. The image firstly is broken into 8x8 blocks of pixels.
2. Then DCT is applied to each block. It works from left to right, top to bottom.
3. Each block is then compressed using quantization table.
4. The array of the compressed blocks which comprise the image is stored in a drastically reduced amount of space.
5. When required, the image is reconstructed through decompression which uses the Inverse Discrete Cosine Transform (IDCT).

ii. Run-Length Encoding (RLE)

This is a lossless algorithm which only furnishes decent compression ratios in specific types of data. It is a form of data compression in which the same data value occurs in many consecutive data elements are stored as single data value along with count. This is most useful on data that contains many such runs. For example, simple graphic images like icons, line drawings, and animations. It may increase the file size as it doesn't have many of the runs, and is not useful

with files. RLE compression can be useful in the following file formats:

- TIFF files
- PDF files

In our case, we are using Huffman algorithm.

IV. THE HUFFMAN ALGORITHM

Huffman coding is an entropy encoding algorithm which is used for lossless data compression in field of computer science and information theory. It refers to the use of a variable-length code table for purpose of encoding a source symbol (such as a character in a file) where the variable-length code table has been derived in a particular manner based on the estimated probability of occurrence for each permissible value of the source symbol. The Huffman coding uses a specific technique for choosing the representation required for each symbol, which results in a prefix-free code (that means, bit string representing some particular symbol is never a prefix of the bit string representing any other symbol) which expresses most common characters using shorter strings of bits which are used for less common source symbols. Huffman was able to design one of the most efficient compression method of this type i.e. no other mapping of individual source symbols to unique strings of bits would produce a smaller average output size when the actual symbol frequencies agree with those used to create the code. A technique was later found to do this in linear time if input probabilities (also known as *weights*) are sorted.

a) Basic Technique

Assume you have a source generating 4 different symbols $\{a1, a2, a3, a4\}$ with probability $\{0.4; 0.35; 0.2; 0.05\}$. Generate a binary tree from left to right taking the two lesser probable symbols, putting them together so as to form another equivalent symbol having probability that equals the sum of two symbols. Keep on doing it until you have just one symbol. After that read the tree backwards from right to the left, assigning different bits to the different branches. The technique works by creating a binary tree of nodes. They can be stored in regular array. The size of array depends on the number of symbols (say N). A node could be either a leaf node or an internal node. Initially, all nodes are leaf nodes, that contain the symbol itself, the weight (frequency of appearance) of the symbol and on optional basis, a link to a parent node which makes it easy to read the code (in reverse) starting from a leaf node. As a common convention, bit '0' represents following the left child and bit '1' represents following the right child. A finished tree has N leaf nodes and N-1 internal nodes.

i. Main Properties

1. Unique Prefix Property i.e. no code is a prefix to any other code.

2. If prior statistics are accurate, then Huffman coding is very good.
 3. The frequencies used can be generic ones for the application domain which are based on average experience.
 4. Huffman coding is optimal when condition of probability of each input symbol is a negative power of two is satisfied.
- ii. *Advantages*
 1. The algorithm is easy to implement.
 2. Produces a lossless compression of images.
 - iii. *Disadvantages*
 1. The efficiency is depended on the accuracy of the statistical model used and type of image.
 2. The algorithm also varies with different formats, but a few get any better than 8:1 compression.

V. FIGURES AND TABLES

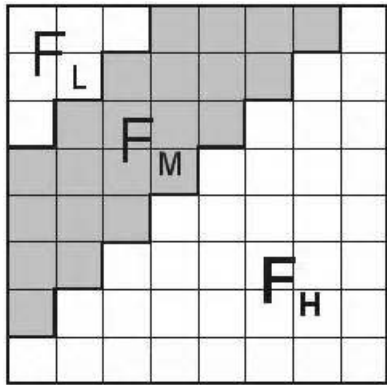


Figure 2.1

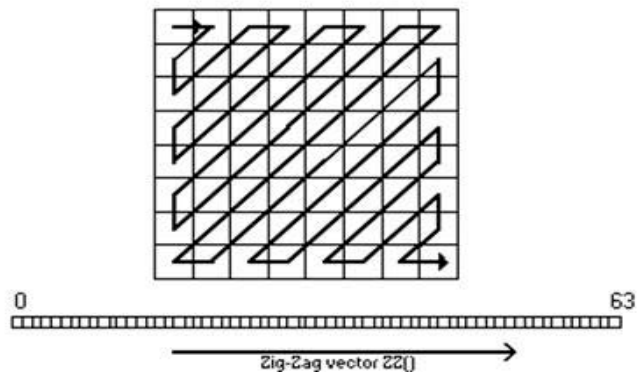


Figure 2.2



Figure 2.3 : Boat



Figure 2.4 : Lena

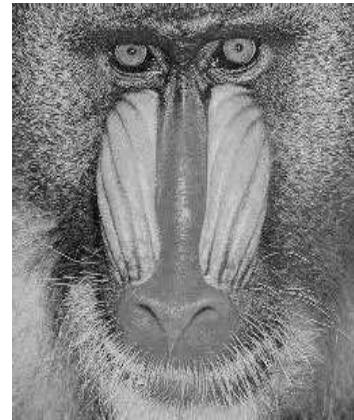


Figure 2.5 : Baboon



Figure 2.6

Table 2.3.1 : Correlation Coefficient (CC) of the Retrieved Watermark at Different Quality Factor

Quality factor(Q)	Baboon	Lena	Boat
5	0.75	0.93	0.90
10	0.775	0.95	0.911
20	0.794	0.99	0.94
40	0.798	0.985	0.9394
45	0.797	0.985	0.9425
50	0.797	0.986	0.9445
60	0.8027	0.9859	0.9468
80	0.972	0.9894	0.9468

VI. RESULT

The watermarked image is subjected to compressions at different quality factors. The watermark from compressed watermarked image is retrieved using the extraction process. The percentage similarity between the extracted watermark and the original watermark is calculated. Fidelity loss of watermarked image is very low. Experimental results show that the proposed scheme is very robust against various image processing operations and geometric attacks. Following are the images which we tested of boat, Lena and Baboon are shown in TABLE 2.3.1.

VII. CONCLUSION

Although the proposed scheme is described for embedding watermark in image, it can be readily adapted for audio watermarking and other forms of watermarking. The study of watermarking technique for digital images shows that its worth exploring the image because of its redundant nature. There is still scope for improvement while working on image watermarking. The results show that this kind of algorithms has a satisfactory performance under image cropping and JPEG lossy compression.

a) Application Area

Watermarks that contain the name of the patient are embedded onto the X-Rays, MRI Scans & other test results which help in instant identification of the result as belonging to a patient and thus avoid mix-ups which can lead to catastrophic consequences. Fig. 2.6 shows an X-ray.

Information like patient name, his age, sex, date of report, medicine before the test, present condition of patient, etc can be embedded in X-ray.

REFERENCES RÉFÉRENCES REFERENCIAS

1. www.en.wikipedia.org/wiki/Watermark
2. www.digimarc.com/technology/about-digital-watermarking.
3. <http://www.watermarks.info/indexi.htm>
4. http://en.wikipedia.org/wiki/Discrete_cosine_transform
5. <http://en.wikipedia.org/wiki/JPEG>
6. "A Feature Based Robust Digital Image Watermarking Scheme" Chih-Wei Tang and Hsueh-Ming Hang, IEEE transactions, vol. 51, April '03.
7. www.imagecompression.info/
8. Watermarking & Stenography <http://www.cl.cam.ac.uk/index.html>



Performance Analysis among Different Classifier Including Naive Bayes, Support Vector Machine and C4.5 for Automatic Weeds Classification

By Md . Mursalin, Md . Motaher Hossain, Md. Kislu Noman
& Md. Shafiul Azam

Pabna University of Science and Technology

Abstract - Weeds are often one of the biggest problems encountered by farmer in conventional agriculture. Maximum productivity of crops can be achieved by proper weeds management. Applying excessive herbicide uniformly throughout the field has an adverse effect on the environment. An automated weed control system which can differentiate the weeds and crops from the digital image could be a feasible solution for this problem. This paper demonstrates Naïve Bayes, SVM (Support Vector Machine) and C 4.5 classification algorithm for classifying the weeds and investigates the performance analysis among these three algorithms. In this study 400 sample images over five species were taken where each and every species contains 80 images. The result has shown that Naïve Bayes classification algorithm achieve the highest accuracy (99.3%) among these three classifier.

Keywords : herbicide, image processing, weed classification, naïve bayes, SVM, C 4.5 classifier.

GJCST-F Classification: 1.5.4



Strictly as per the compliance and regulations of:



Performance Analysis among Different Classifier Including Naive Bayes, Support Vector Machine and C4.5 for Automatic Weeds Classification

Md. Mursalin^α, Md. Motaher Hossain^σ, Md. Kislu Noman^ρ & Md. Shafiul Azam^ω

Abstract - Weeds are often one of the biggest problems encountered by farmer in conventional agriculture. Maximum productivity of crops can be achieved by proper weeds management. Applying excessive herbicide uniformly throughout the field has an adverse effect on the environment. An automated weed control system which can differentiate the weeds and crops from the digital image could be a feasible solution for this problem. This paper demonstrates Naïve Bayes, SVM (Support Vector Machine) and C 4.5 classification algorithm for classifying the weeds and investigates the performance analysis among these three algorithms. In this study 400 sample images over five species were taken where each and every species contains 80 images. The result has shown that Naïve Bayes classification algorithm achieve the highest accuracy (99.3%) among these three classifier.

Keywords : herbicide, image processing, weed classification, naïve bayes, SVM, C 4.5 classifier.

1. INTRODUCTION

Weeds are those unwanted plants which grow in the area not belongs to them and cause more negative impact on economy income. It competes with the crop for resource such as soil, water, sunlight and fertilizer. So the production efficiency and quality of economic crops would decrease when the weeds are out of control. Statistics has been shown that the worldwide estimated potential loss due to all kinds of pests was at 40%-80%; besides them the potential losses for weeds were found 34% which is the highest of all pests [8]. As a result, better weed control system must be deployed to sustain the productivity without hampering the environment. Currently several weed control policies exist e.g. removing weeds manually by human laborers, crop rotation, mechanical cultivation, and chemical herbicides.

The huge rate of herbicide in the world causes negative impact on the ecological environment and the survival of species. It has also arises some economic concern. In year of 2005, the total estimated cost of applied herbicides was \$16 billion in United States [8].

Author ^α ρ : Lecturer, Department of Computer Science and Engineering, Pabna University of Science and Technology.
E-mails : mursalin.mithu@gmail.com, md.k.noman@gmail.com

Author ^σ : Lecturer, Department of Electrical and Electronic Engineering, Pabna University of Science and Technology.
E-mail : motahertex@yahoo.com

Author ^ω : Assistant Professor, Department of Computer Science and Engineering, Pabna University of Science and Technology.
E-mail : shahincseru@gmail.com

The main cost ineffective and strategic problems in herbicides system is that they are applied on the field uniformly. Generally the volumes of weeds are more in some specific region but herbicides are still applied regardless. In addition, applying herbicides manually is very costly and time consuming. Statistics has been shown that if same kinds of herbicides are applied repeatedly for reducing the unwanted weeds then there is a good possibility that they become tolerant to those types of herbicides [6]. Moreover some herbicides contaminate the ground water even though it applied in the soil. Thus farmers need more sophisticated alternative weed control techniques which will reduce the usages of chemicals and provide safety for the overall ecosystem.

Several researches have been done for investigating fruitful solution for controlling the weeds without collapsing down the environment. The machine vision technique has the ability to differentiate the crops from weeds so that herbicides can be applied effectively. In this technique image are captured by a digital camera from different parts of a crop field so that weeds can be identified properly. Shearer, et al. [10] has developed a photo sensor plant detection system which has the ability of detecting and spraying only the green plant. Jiang Zhengrong, et al [7] has proposed automatic weed identification based on image processing technology. They have investigated the spectrum analysis, color identification, texture assessment for weed identification. In [3], weeds and crops are classified by SVM and achieved 98% where using Bayesian classifier achieved 95% accuracy over 224 test images. Weis et al, proposed a sensor related analysis techniques of weed detection system [14]. In [12], comparison of different classification algorithms has been shown for weed detection based on shape features. For selective herbicide application a model has been proposed [1] with 95% accuracy, which categorizes images into narrow and broad classes based on the Histogram Maxima using a thresholding technique. In [13] calculation of various shape features for identifying weeds in digital images has been shown. Active shape models can identify young weed seedlings with the accuracy of 65% to 90% [11].

The aim of this paper is to introduce a model which will classify weeds and crops from digital images

using Naïve Bayes, SVM and C 4.5 classifier and to evaluate their performance in an automated weed control system. These three classification technique are examined to find the optimum solution. Naïve Bayes classifier has been preferred as it is simple, effective and fast among other classification algorithms.

II. METHODS

a) Process Flow

The overall procedure of this paper has shown in Fig: 1. Images were captured by a digital camera with 4.9-24.5 mm lens. The position of camera was 90 degree angle form the ground that means vertically towards the ground. The distance between camera lens and the ground was 1.3 feet. Photo shed was used for keeping same light intensity. 1024x768 photo resolution was set for capturing the color image of weeds and crops. All images were taken from the capsicum filed. There were five species including the capsicum. Other four species were considered as weed for the capsicum.

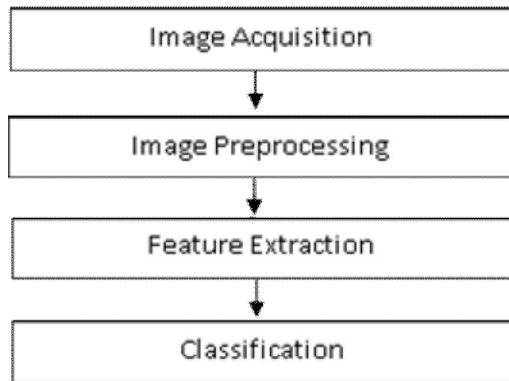


Figure 1 : Step by step procedure

b) Image Preprocessing

The high resolution image was converted to 225x175 pixels in order to minimize the computation cost. Color segmentation based image-processing has been done for distinguishing plants from background objects where objects are one of two classes as plants and background. The plants in the field images must be properly segmented otherwise extraction of features will give improper results. Each and every pixel of rgb (red, green, blue) image an exhibitor value 'e' was calculated by using color component for enhancing the green plant in compare to the background:

$$e = 2 \times g - r - b \quad (1)$$

The rgb color plant images were converted to grey images after calculation of e value. Binarization technique with global threshold was performed on these images to separate plants form the background. Composite Laplacian mask was used for further enhancement of the grey-scale image [3]. As the sharpening procedure is sensitive to noise, a linear smoothing method known as median filter was used which successfully reduce impulse noise [4]. Otsu's

method [9] an effective technique was used to select the proper binarization threshold value. If the pixel value 'p' is smaller than threshold value't' were referred as soil in the binary image. In binary image '0' indicates the background where '1' indicates the plant.

To remove the noise from binary image, at first morphological opening has applied. In this operation, an erosion operation is followed by a dilation operation. It makes smooth the image by eliminating small pixel regions. The erosion and dilation were combined in reverse order for morphological close operation. This close operation fills the small holes in image [5].

c) Features Extraction

Ten features were extracted from the binary images (Fig 2). These features were decomposed as shape, color and moment invariants. The shape features were divided into two categories: size dependent and size independent. Size dependent descriptors are area, perimeter, thickness, convex area and convex perimeter. The size dependent features were combined to present size dependent shape features:

$$formfactor = 4 \times \pi \times \frac{area}{perimeter^2} \quad (2)$$

$$elongatedness = \frac{area}{thickness^2} \quad (3)$$

$$convexity = \frac{convex_perimeter}{perimeter} \quad (4)$$

$$solidity = \frac{area}{convex_area} \quad (5)$$

For making the color features consistent with various lighting conditions, each and every color component was divided by sum of all the three color components. Here the color features were mean value and the standard deviation. The scope of an object area is measured by moment invariant (1 N, 2 N, 3 N, 4 N) [2] which consists of geometric transformation such as scaling, translation and rotation. Here in this study only central moments are considered.

d) Classification using Naïve Bayes Classifier

The Naïve Bayes classifier is a simple but effective classifier has been used here to minimize the computation cost. Let $\vec{a}$ be a vector we want to classify and c be a possible class. Using Bayes formula first we transform the probability $P(c | \vec{a})$:

$$P(c | \vec{a}) = P(c) \times \frac{P(\vec{a} | c)}{P(\vec{a})} \quad (6)$$

$P(c)$ can be estimated from training data. Considering the conditional independence of the elements of a vector $P(\vec{a} | c)$ is decomposed as follows,

$$P(\vec{a} | c) = \prod_{i=1}^D P(a_i | c) \quad (7)$$

Where, a_i is the i^{th} element of $\vec{a}$. Now, combining both equations we get:

$$P(c | \vec{a}) = P(c) \times \frac{\prod_{i=1}^D P(a_i | c)}{P(\vec{a})} \quad (8)$$

From this final equation we can calculate $P(c | \vec{a})$ and classify $\vec{a}$ into the class with the highest $P(c | \vec{a})$.

A classification process in Naïve Bayes classifier requires first train the classifier using labeled data. Then classify unlabeled examples with assigning probabilistic labels to them. In this paper we consider binary classification as weed and crops.

Let k_i be the probabilistic label of i^{th} example illustrate the probability that it belongs to weed class. If the proportion of weed class examples in unlabeled data is different from labeled data then the probabilistic labels were calibrated. The main theme of the calibration is to shift all the probability values of unlabeled data to the scope that the class distribution of unlabeled data becomes alike to that of labeled data. In [15] the whole calibration process has shown.

e) Classification using C4.5 Classifier

Using the concept of information entropy, [18] C4.5 builds decision trees from a set of training data, in the same way as ID3 (Iterative Dichotomiser 3). [16] The training data is a set $S = (s_1, s_2, \dots, s_n)$ of already classified samples. Each sample $s_i = (x_1, x_2, \dots, x_n)$ is a vector where x_i represent attributes or features of the sample. The training data is augmented with a vector $C = c_1, c_2, \dots, c_n$ where c_i represent the class to which each sample belongs. [17] C4.5 algorithm selects the attribute of the data that most effectively splits its set of samples into subsets enriched in one class or the other at each node of the tree. The splitting criterion is the normalized information gain (difference in entropy). The attribute with the highest normalized information gain is chosen to make the decision. The C4.5 algorithm then recurses on the smaller sublists. Base case of this algorithm:

- All the samples in the list belong to the same class. When this happens, it simply creates a leaf node for the decision tree saying to choose that class.
- None of the features provide any information gain. In this case, C4.5 creates a decision node higher up the tree using the expected value of the class.
- Instance of previously-unseen class encountered. Again, C4.5 creates a decision node higher up the tree using the expected value.

f) Classification using SVM

First task of SVM classifier requires separating the dataset into two different parts. First one is used for training and second one is used for testing. A class label and the corresponding image features have been assigned to each instance in the training set. When the features values are provided, SVM generates a classification model which is used to predict the class labels of the test data depending on training data. Each instance is represented by an n -dimensional feature vector, $V = (v_1, v_2, \dots, v_n)$. Here, 'V' depicts n measurements made on an instance of n features. The dataset is normalized before use because the feature values for the dataset can have ranges that vary in scale. The LIBSVM 2.91 [19] library was used to implement the support vector classification where each feature value of the dataset was scaled to the range of [0, 1]. The RBF (Radial-Basis Function) kernel was used for both SVM training and testing which mapped samples nonlinearly onto a higher dimensional space. For this reason, this kernel is able to handle cases where nonlinear relationship exists between class labels and features. A commonly used radial basis function [3] is:

$$K(v_i, v_j) = \exp(-\gamma \|v_i - v_j\|^2), \gamma > 0 \quad (9)$$

Where,

$$\|v_i - v_j\|^2 = (v_i - v_j)^t (v_i - v_j) \quad (10)$$

Here, ' v_i ' and ' v_j ' are n -dimensional feature vectors. Implementation of the RBF kernel in LIBSVM 2.91 requires two parameters: ' γ ' and a penalization parameter, 'C' [19]. Appropriate values of 'C' and ' γ ' should be specified to achieve a high accuracy rate in classification. By repeated experiments [3], $C = 1.00$ and $\gamma = 1/n$ were chosen.

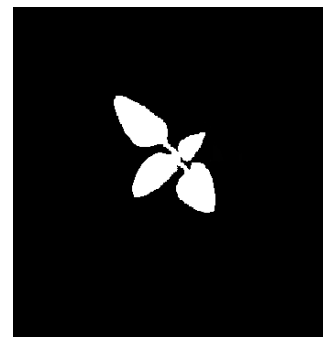


Figure 2 : Binary image after open-close operation

III. RESULT AND DISCUSSION

In this paper, each setting is evaluated by using 10-fold cross-validation procedure. 10-fold cross-validation procedure needs portioning the whole training set into 10 subgroups. Each and every subgroup has an equal number of instances. In this training process, one subgroup is tested with the remaining nine subgroups.

As a result, over fitting protection is ensured and smooth outcome for the actual computing is achieved.

Three hundred sample image data were trained and one hundred sample image data were tested. The result of 10-fold cross-validation of Naïve Bayes classifier using ten features was found 99.3% accurate. 98.24% and 97.86% accuracy has been achieved using 10-fold cross-validation of SVM and C4.5 classifier consecutively. Table 1 has shown the success rate comparison using all features. The number of features has been reduced to minimize the computational complexity. This study has experimented on fifteen features and by using forward-selection and backward-elimination methods 10 features achieved the optimum accuracy rate. Selected features were convexity, mean value of 'r', mean value of 'b', standard deviation of 'r', standard deviation of 'b', $\ln(N_1)$ of area, $\ln(N_2)$ of area, $\ln(N_3)$ of area, $\ln(N_4)$ of area, $\ln(N_2)$ of perimeter.

In present study the capsicum, cogon grass and marsh herb were successfully classified. Other two species had some misclassifications. Table 2 shows the comparative accuracy rate for each species. Here each and every species has trained with 60 samples and 20 sample images were used for testing whether the classifier can successfully classify or not.

Method	Average Success Rate(%)
Naïve Bayes	95.4
SVM	95
C4.5	91.6

Table 1 : Classification result using all features

Plants Name	Naïve Bayes (%)	SVM (%)	C 4.5 (%)
Capsicum	100	100	100
Burcucumber	98.5	96.2	94.3
Cogongrass	100	100	100
Marsh herb	100	100	99
Pigweed	98	95	96
Average Accuracy Rate	99.3	98.24	97.86

Table 2 : Classification result using set of best features

IV. CONCLUSION

Our main goal of this work is to find a solution which will minimize the operating cost as well as maximize the result. In this paper, three different classifier including Naïve Bayes, SVM and C4.5 have been evaluated to classify the weeds and crops. Compare to SVM and C4.5, Naïve Bayes classifier

obtains highest result. The future work will focus on wavelet transformation in image preprocessing steps. We will also study the optimization technique for these classifiers and ensure that the large training set will not cause over fitting problem.

V. ACKNOWLEDGEMENT

With due homage and honor we wish to express our gratitude to Almighty. We wish to express our love and gratitude to our beloved families; for their understanding & endless love, through the duration of studies. We are grateful for their continuous support and help, which taught us to be more responsible for own good and for others.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Irshad Ahmad, Abdul Muhamin Naeem, Muhammad Islam, Shahid Nawaz, "Weed Classification using Histogram Maxima with Threshold for Selective Herbicide Applications", In the Proceedings of the 19th International Conference on Computer, Information and Systems Science and Engineering (CISE 2007), Bangkok (Thailand), 19: 2007, 331-334.
2. Faisal Ahmed, Hawlader Abdullah Al-Mamun, A.S.M Hossain, Emam Hossain, Paul Kwan, "Classification of crops and weeds from digital images: A support vector machine approach", Crop Protection, Volume 40, ISSN 0261-2194, 10.1016/j.cropro.2012.04.024, October 2012, Pages 98-104.
3. Faisal Ahmed, A.S.M Hossain, Emam Hossain, Hawlader Abdullah Al-Mamun, Paul Kwan, "Performance Analysis of Support Vector Machine and Bayesian Classifier for Crop and Weed Classification from Digital Images", World Applied Sciences Journal, Volume 12 (4), ISSN 1818-4952, IOS Publications, 2011, Pages: 432-440.
4. Fu, J.C., J.W. Chai, S.T.C. Wong, J.J. Deng, J.Y. Yeh, "De-noising of left ventricular myocardial borders in magnetic resonance images", Magnetic Resonance Imaging, Volume 20(9), 2002, Pages: 649-657.
5. Rafael C. González, Richard Eugene Woods "Digital Image Processing", Pearson Education, Singapore, 2nd edition, pp: 528, 2004.
6. "International survey of herbicide tolerant weeds", <http://www.weedscience.org/in.asp> (accessed on 2, June, 2012).
7. Jiang Zhengrong, Li Zhengxi, "Automatic Weed Identification Based on Image Processing Technology", the 3rd International Conference on Machine Vision, ISBN: C 978-1-4244-8889-6, 2011.
8. E. -C. Oerke, "Crop losses to pests" The Journal of Agricultural Science and Volume 144(1), doi: 10.1017/S0021859605005708, 2006, pp: 31-43.

9. Otsu, "A Threshold Selection Method from Gray-Level Histograms", IEEE Transactions on Systems, Man and Cybernetics, SMC-9(1): 62- 66, 1979.
10. Shearer, S.A. and P.T. Jones, et al., "Selective Application of Post-Emergence Herbicides Using Photoelectric", ASABE Technical Library, Transactions of the ASAE, Volume 34(4), 1991, pp: 1661-1666.
11. Sogaard, H.T., "Weed Classification by Active Shape Models", Biosystems Engineering, Volume 91(3), 2005, pp: 271-281.
12. Weis, M., T. Rumpf, R. Gerhards, L. Plumer, "Comparison of different classification algorithms for weed detection from images based on shape parameters", In the Proceedings of the 1st International Workshop on Computer Image Analysis in Agriculture, Potsdam, Germany, 2009, Volume 53, pp : 460-463.
13. Weis, M., R. Gerhards, "Feature extraction for the identification of weed species in digital images for the purpose of site-specific weed control", In the Proceedings of the 6th European Conference on Precision Agriculture, Skiathos, Greece, Volume 6, 2007, pp: 537-544.
14. Weis, M., C. Gutjahr, V. Rueda, R. Gerhards, C. Ritter, F. Scholderle, "Precision Crop Protection - the Challenge and Use of Heterogeneity", 1st Edition, Chapter Detection and identification of weeds, Springer Dordrecht Heidelberg London, New York, pp: 119-134, 2010.
15. Yoshimasa Tsuruoka and Jun'ichi Tsujii, "Training a Naive Bayes Classifier via the EM Algorithm with a Class Distribution Constraint", www.acl.ldc.upenn.edu/W/W03/W03-0417.pdf.
16. S. Anupama Kumar, Dr. Vijayalaskshmi, "Implication of Classification Techniques in Predicting Students Recital", International Journal of Data Mining & Knowledge Management Process (IJDMP) vol.1, No.5, September 2011.
17. Sonal Athavale, Neelabh Sao, "Classification on Moving Object Trajectories", International Journal of Advanced Technology & Engineering Research (IJATER), ISSN No: 2250-3536, vol 2, Issue 2, May 2012.
18. Quinlan, J. R. C4.5: Programs for Machine Learning. Morgan Kaufmann Publishers, 1993.
19. LIBSVM-A Library for Support Vector Machines, <http://www.csie.ntu.edu.tw/~cjlin/libsvm/> (accessed on 10 June, 2012).



This page is intentionally left blank



Software for Drawing Jittered Plot: A Graphical Counterpart of Anova

By Dr. Dibyo Jyoti Bhattacharjee & Dr. Rupak Gupta

Assam University, Silchar, Assam & University of Technology & Management, Shillong, Meghalaya

Abstract - The Analysis of Variance is a robust tool for comparing significant difference between means of multiple series. The paper discusses an existing graphical tool that can be considered as a visual counterpart for the same purpose, and proposes some extensions to the plot. The jittered plot for different factors each at two levels along with their means and corresponding confidence intervals is the maximum extension that is been discussed in the paper. Software is also developed using Visual Basic 6.0 that can be used for drawing the jittered plot along with all the discussed extensions. The plot drawn in the software is compatible with other word processors. Thus the reader can use his/her's own data set and obtain the corresponding plot.

Keywords : *graphical data analysis, anova, jittered plot, multiple comparison.*

GJCST-F Classification: *1.5.m*



Strictly as per the compliance and regulations of:



Software for Drawing Jittered Plot: A Graphical Counterpart of Anova

Dr. Dibyo Jyoti Bhattacharjee^a & Dr. Rupak Gupta^o

Abstract - The Analysis of Variance is a robust tool for comparing significant difference between means of multiple series. The paper discusses an existing graphical tool that can be considered as a visual counterpart for the same purpose, and proposes some extensions to the plot. The jittered plot for different factors each at two levels along with their means and corresponding confidence intervals is the maximum extension that is been discussed in the paper. Software is also developed using Visual Basic 6.0 that can be used for drawing the jittered plot along with all the discussed extensions. The plot drawn in the software is compatible with other word processors. Thus the reader can use his/her's own data set and obtain the corresponding plot.

Keywords : graphical data analysis, anova, jittered plot, multiple comparison.

I. INTRODUCTION

A graphical complement to ANOVA provides us an opportunity to visualize how the central values of several groups vary amongst themselves and also how the values are spread over in different groups. These provide more information than a statistical test of significance. An ANOVA test can only measure if there is a significant difference between means and conceals information related to the magnitude of the difference but a graph can show both, in addition to that it also helps in the visualization of the scattering of data across the different samples.

II. DESCRIPTION OF THE PLOT

A jittered plot is commonly used to represent multiple groups in a graph. The groups are taken along the X axis which are been separated by a distance. The groups may be named or some identifying number can be used. Above each of the identifier variables (along Y axis) are plotted in the form of small circles the values of the variables. The circles are kept transparent in order to avoid overlapping. Another simple technique of avoiding overlapping may be to generate a random number between $[-0.2, 0.2]$ and add it to the group identifier variable before plotting. This disturbs the horizontal alignment of the points belonging to a particular group and hence acts as a solution to the problem of overlapping in a group. The mean (or median) of the series is denoted by a small but thick straight line within

the cluster of points in each group. This makes the means readily visible. However, one can also view the spread of the data around the central value, range, symmetry around the central point and so on for each of the groups. Hence, the plot can be used as a tool for comparison of central values, deviation, range, concentration around central value and so on for multiple data sets.

III. A SAMPLE DATA SET AND PLOT

For drawing the jittered plot the following data is considered originally from Snedecor (1956) which gives the birthweights of Poland China pigs for four litters in pounds.

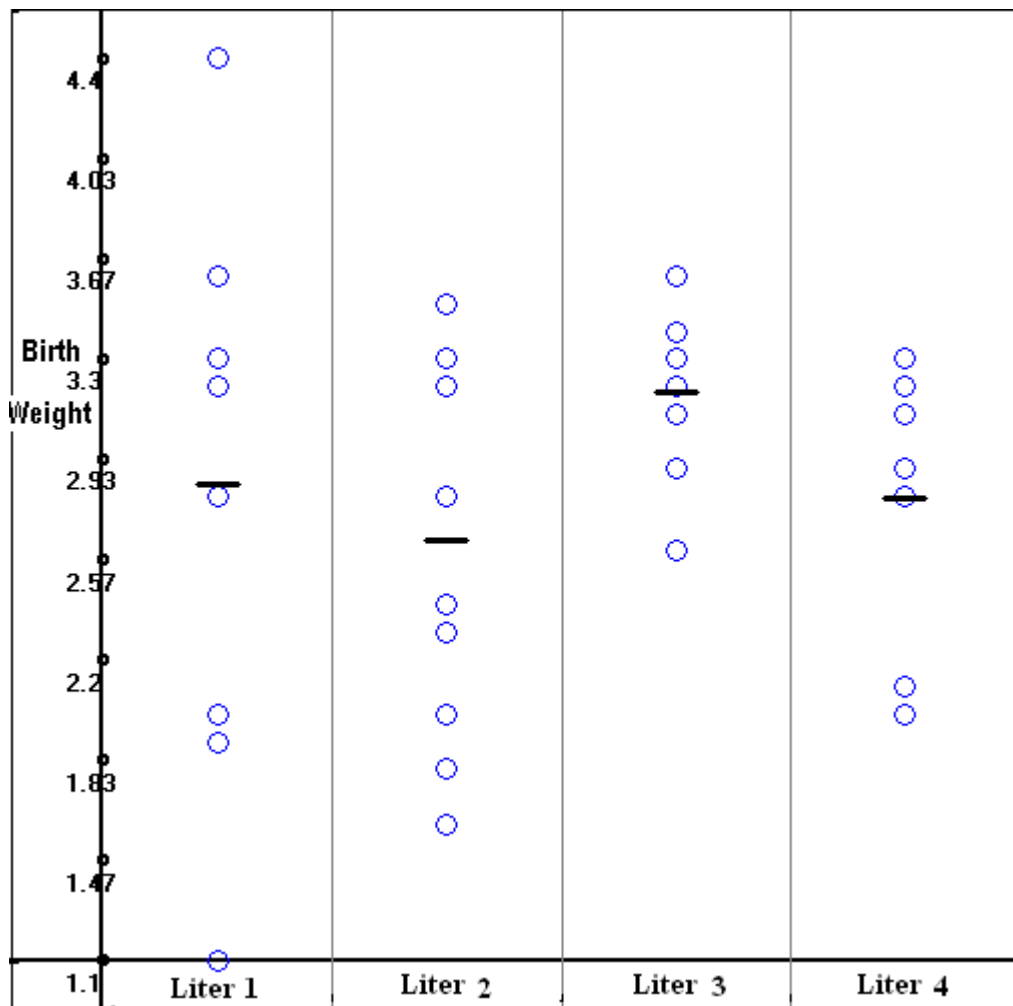
Table 1 : Birthweight of Poland China Pigs in Pounds

Liter 1	Liter 2	Liter 3	Liter 4
2.0	3.5	3.3	3.2
2.8	2.8	3.6	3.3
3.3	3.2	2.6	3.2
3.2	3.5	3.1	2.9
4.4	2.3	3.2	2.0
3.6	2.4	3.3	2.0
1.9	2.0	2.9	2.1
3.3	1.6	3.4	3.1
2.8	1.8	3.2	2.8
1.1	3.3	3.2	3.3

Author a : Department of Business Administration Assam University, Silchar, Assam.

Author o : Department of Mathematics, University of Technology & Management Shillong, Meghalaya.

Figure 1 : A jittered plot for data in Table 1

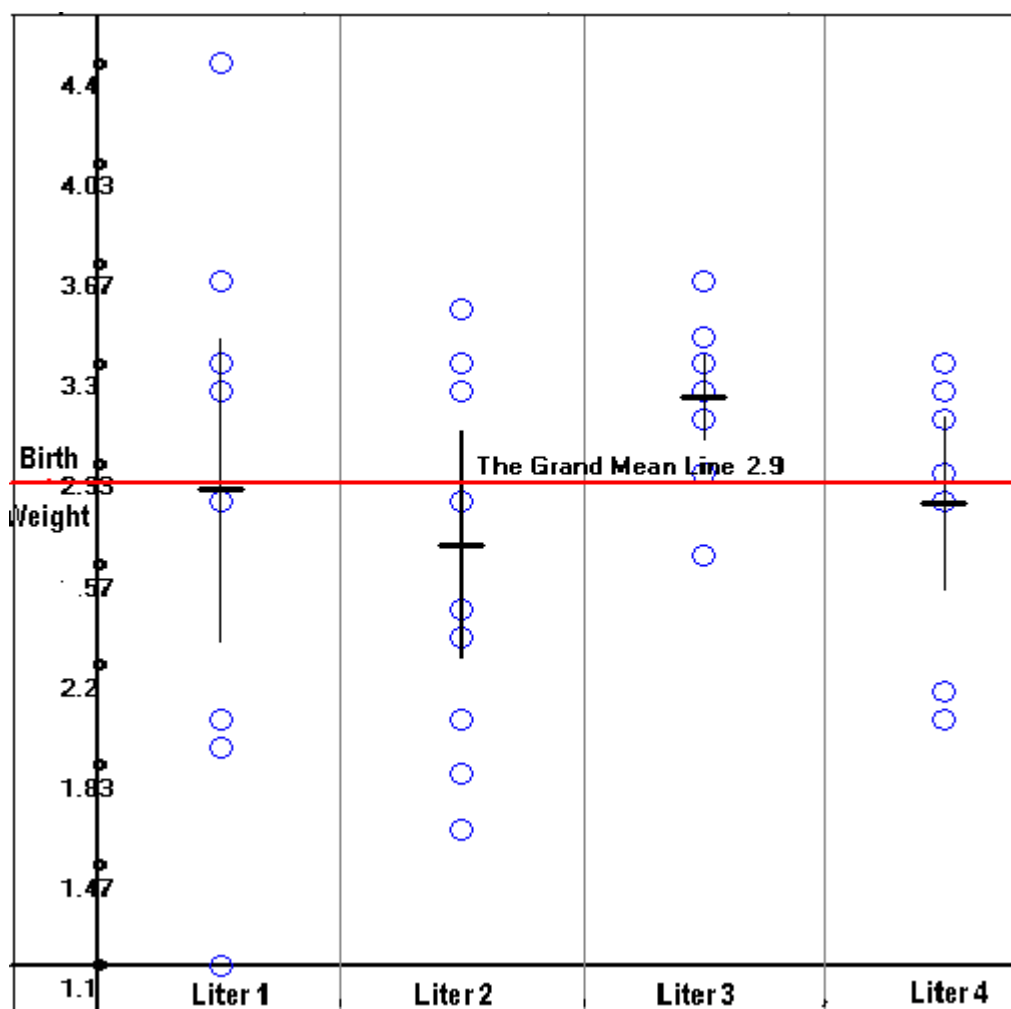


From the jittered plot drawn above one can understand that the means of the second group is less, compared to that of the other groups. Also the spread of the data is seen to be more in the first two groups and minimum in the third one.

IV. SOME EXTENSIONS TO THE PLOT

In order to understand the difference of means of any two pairs, it is reasonable to find the confidence interval for the means for a given level of significance and accordingly plot them in the jittered plot. The confidence interval of the mean of each group is represented by two straight lines drawn from the thick, small line used to represent the mean of each series. One line is drawn from the mean to the upper confidence limit and the other one is drawn from the mean to the lower confidence limit. This is repeated for each of the groups. Thus we get a confidence line in each group with the group mean exactly at the center. The grand mean may also be computed and may be represented by a line parallel to x axis running across the plot with appropriate ordinate. Figure 2, shows the extended plot for the Poland China Pig data provided in Table 1.

Figure 2 : A jittered plot with proposed confidence intervals data in Table 1



From the diagram it is clear that the mean of the second group with that of the third differs significantly as can be visualized from the difference in the vertical alignment of the confidence bands of the two groups. Based on the same logic, we can conclude that there is no significant difference between the means of the first two series.

Now let us move a step forward and think of groups each having two levels. Thus here one must be interested to note the impact of various levels in each of the groups. The same jittered plot may be extended

further to visualize both the levels. This allows comparison between levels within a group as well as between groups. The data given below is taken from Barlett (1936).

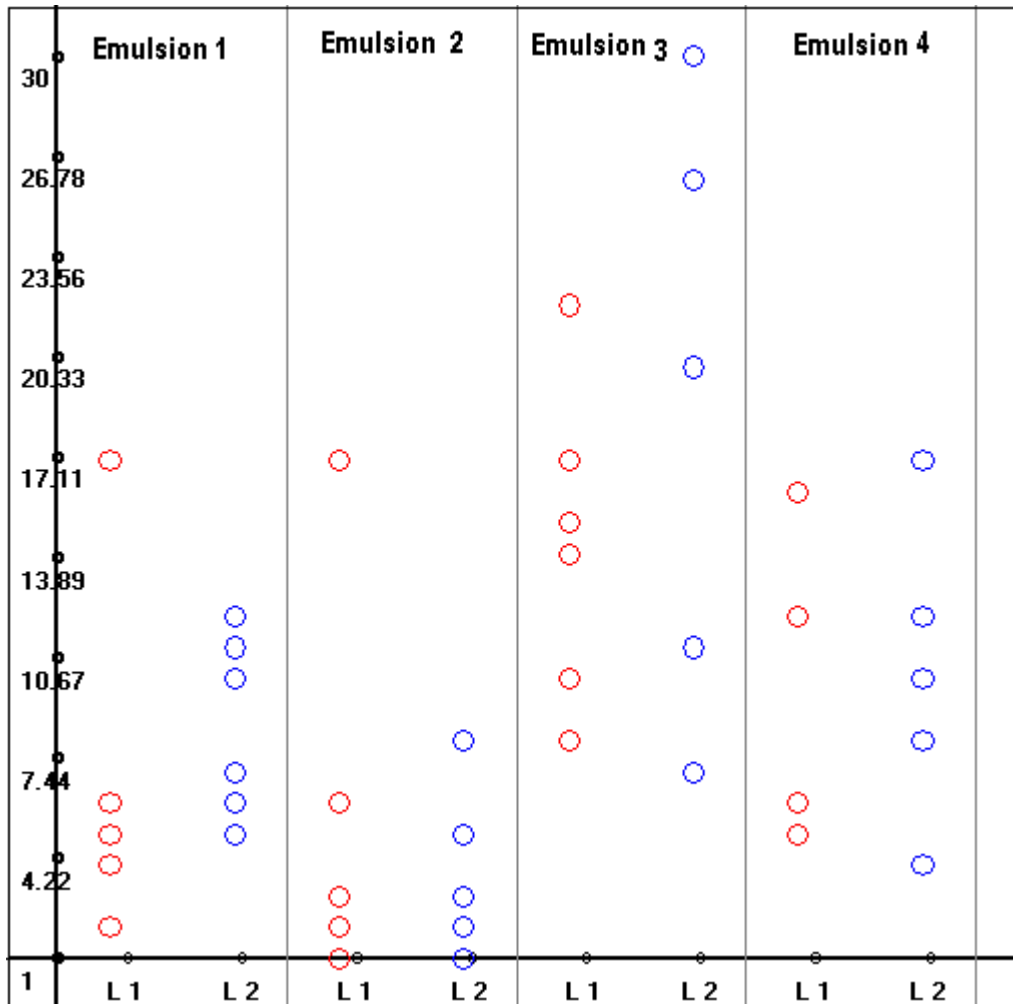
Table 2 : Counts of surviving latherjackets for different Controls and Emulsions each at two levels

Treatment		Emulsion 1	Emulsion 2	Emulsion 3	Emulsion 4
Block					
	Level 1				
1	Level 1	6	17	8	12
	Level 2	10	8	11	17
2	Level 1	4	3	15	6
	Level 2	7	2	20	4
3	Level 1	4	6	10	12
	Level 2	12	3	7	10
4	Level 1	5	1	17	5
	Level 2	5	1	26	8
5	Level 1	2	2	14	12
	Level 2	6	5	11	12
6	Level 1	17	6	22	16
	Level 2	11	5	30	4

The figure below depicts the data through a jittered plot. Here observations at different levels, under

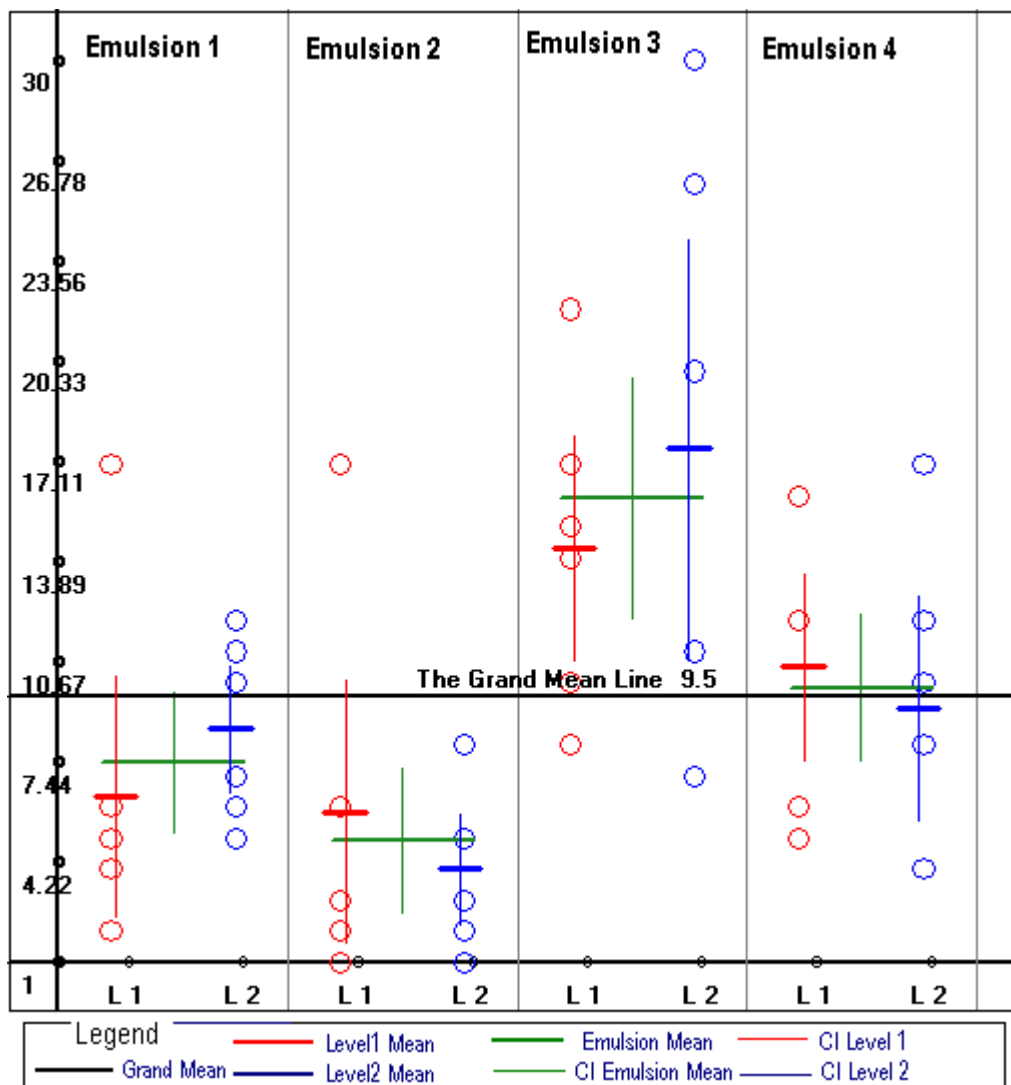
the same Emulsion are been aligned differently or may be symbolized differently or both.

Figure 3 : Extended jittered plot with each factor at two levels for latherjackets data



An extension to this plot obviously follows, with group means, treatment means along with the confidence bands. The different means and confidence bands are plotted in the fashion discussed earlier with appropriate color. Figure 4, produces such a plot for latherjacket data with the discussed extensions.

Figure 4 : A Jittered Plot with confidence bands for data at two levels for the latherjacket data



The legend probably explains all. Here one can compare the performance of the same level at different factors as well as the performance of a factor at different levels. In addition to this the confidence intervals for different means helps one to get an idea about the significant difference between any pair of means. The position of any type of mean can be visualized with respect to the grand mean.

To sum up it may be said that the simple plot is actually a multipurpose plot and it can be used for the following purpose:

- Location and scale difference of the response variable under different treatments.
- Location and scale difference of the response variable at different levels for the same treatment.
- Presence of outliers.
- Comparison of treatment means amongst them as well as with the grand mean.
- Spread of treatment means around the grand mean.
- Comparison of means at different levels for a given treatment.
- Comparison of means at same level for different treatments.

V. SOFTWARE DEVELOPMENT

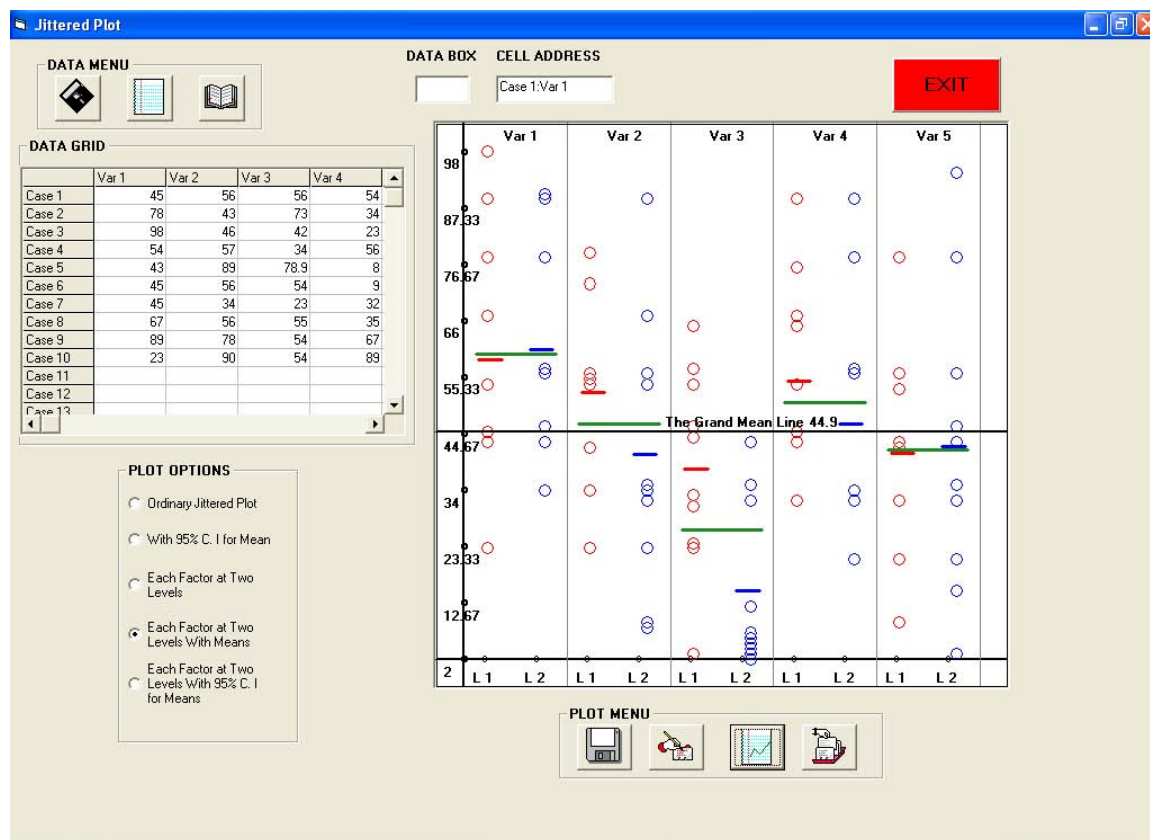
It is always said that the graphical techniques are simpler compared to the other statistical tools. But this is not always the fact. As in this case, we find that a lot of computation is required before one actually starts drawing the plot. Thus the authors think that an appropriate way is to develop stand alone software for drawing the jittered plot with all the proposed extensions and accordingly help the reader to load it in his computer and then draw the appropriate figure. Accordingly stand alone software has been developed using Visual Basic 6.0 the code for which is not provided.

The software can be loaded very simply in the usual manner by clicking the "Setup" present in the

Jittered Plot folder. On going through the commonly used procedure the software gets loaded in the location desired in the computer. The software can be accessed

from the Programs list under the start menu. On clicking there the software opens which looks like the one in the figure below.

Figure 5: The opening screen of the Software

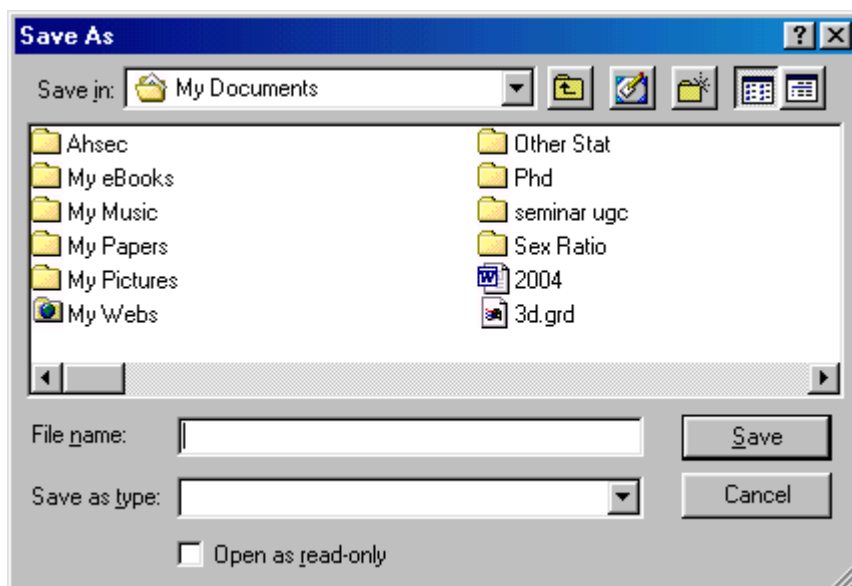


a) Using the Software

For entering data into the software the "Data Box" is used. The value is entered in the "Data Box" and then the return key is pressed. The value in the data box is transferred to the first cell under the first column in the "Data Grid". The next cell in the same column now gets ready to accept the data. The cursor however remains in the data box, where the data is entered and the return key is pressed. The cell address shows the current cell number i.e. the cell that remains in the activated stage. It must be remembered that each variable should have equal number of observations. The data grid can space a maximum of 10 variables and 30 cases under each variable. If the user wants to type data in any other cell then the default cell, just click that cell and accordingly type the data in the data box.

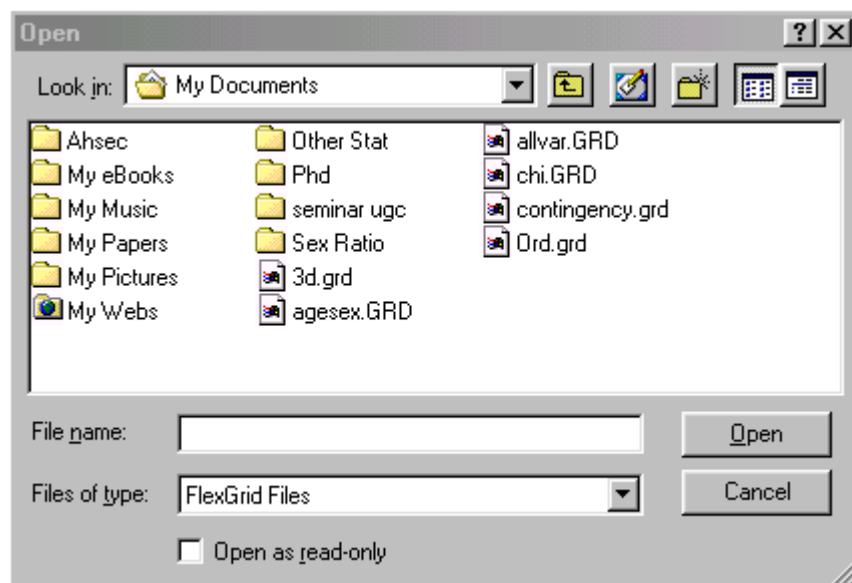
To save the data, click the save option in the "Data Menu" and the save dialogue box appears. Accordingly the following dialog will be displayed in the screen where you can type the desired file name and click on 'Save'. If one wants to save the data in a different folder (directory) then select the folder name and then type in the file name and click 'Save'. The file will be saved with the extension '.grd'.

Figure 6 : The 'Save As' Dialog Box



In order to open a particular data that has already been saved and stored in a particular folder, click on 'Open' button of the "Data Menu" and accordingly producing the following dialog box.

Figure 7 : The 'Open' Dialog Box



Accordingly, the required file that opens would go to the data grid of the software. Using the dialog box one may even search the other files of '.grd' extension, if the required files are in some other folder. In order to create a new data set one can click 'New' option in the "Data Menu". Accordingly, any opened file gets closed and the data grid gets refreshed.

Once the data has been entered properly, next the appropriate plot can be chosen from "Plot Options". Now on clicking the draw button in the "Plot Menu" the desired jittered plot is created in the space provided just above the menu. The "Plot Menu" also provides option to copy the plot to clipboard and to paste in some other appropriate software. The save option in the menu also

provides option to save the plot as file with 'bmp' extension in any folder.

VI. CONCLUSION

The plot discussed above with all its extension may be of specific interest to social scientists and researcher who are not comfortable with ANOVA. The plot can also be used with ANOVA in order to visualize the difference between the means as well as many more features of the data set as discussed earlier. A new plot should be supported with software that can help the reader to draw the plot for user defined data set. So, one has to take the support of a program in high level language with a hardware support of moderate

resolution computer graphic system to develop such a package.

REFERENCES RÉFÉRENCES REFERENCIAS

1. **Barlett , M.S. (1936)** Some notes on insecticide tests in the laboratory and in the field. Journal of the Royal Statistical Society, Supplement, 3, 185–194.
2. **Snedecor, G.W. (1956)** Statistical Methods, 5th Edition, Ames: Iowa State College Press.





Data Adaptive Multi-Band Approach for Robust Audio Watermarking using Adaptive Threshold

By K.M. Ibrahim Khalilullah, A K M Wasif Abedin, Mahfuza Khatun,
Mohammad Tareq & Md. A.R. Pramanik

Dhaka International University, Dhaka, Bangladesh

Abstract - This paper presents a data adaptive digital audio watermarking process with the use of Empirical Mode Decomposition (EMD) and Hilbert Transform (HT). The host audio signal is decomposed into several Intrinsic Mode Functions (IMFs). A binary image or a set of -1 and 1, which is obtained by mapping from standard normal distributed pseudo random numbers, is embedded as secret information into the significant coefficients of the highest energetic IMF that are greater than a specified adaptive threshold. Thus watermarked IMF is less sensitive to common signal processing attack. As a result this method increases more robustness. The experimental results show that the proposed method has good imperceptibility and robustness under common signal processing attacks such as additive noise (in time domain and in frequency domain), low pass filtering, re-sampling, re-quantization, MP3 compression, and sound processing effects such as, delay, a natural sounding reverberation (Schroeder's Reverberator), flanging effect, equalization effect, etc. A Bit Error calculation equation is provided to measure the accuracy of the extracted watermark.

Keywords : audio watermarking, empirical mode decomposition, hilbert transform, wavelet transform, adaptive threshold, binary image.

GJCST-F Classification: 1.4.5



DATA ADAPTIVE MULTI-BAND APPROACH FOR ROBUST AUDIO WATERMARKING USING ADAPTIVE THRESHOLD

Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

Data Adaptive Multi-Band Approach for Robust Audio Watermarking using Adaptive Threshold

K.M. Ibrahim Khalilullah^α, A K M Wasif Abedin^σ, Mahfuza Khatun^ρ, Mohammad Tareq^ω
& Md. A.R. Pramanik[¥]

Abstract - This paper presents a data adaptive digital audio watermarking process with the use of Empirical Mode Decomposition (EMD) and Hilbert Transform (HT). The host audio signal is decomposed into several Intrinsic Mode Functions (IMFs). A binary image or a set of -1 and 1, which is obtained by mapping from standard normal distributed pseudo random numbers, is embedded as secret information into the significant coefficients of the highest energetic IMF that are greater than a specified adaptive threshold. Thus watermarked IMF is less sensitive to common signal processing attack. As a result this method increases more robustness. The experimental results show that the proposed method has good imperceptibility and robustness under common signal processing attacks such as additive noise (in time domain and in frequency domain), low pass filtering, re-sampling, re-quantization, MP3 compression, and sound processing effects such as, delay, a natural sounding reverberation (Schroeder's Reverberator), flanging effect, equalization effect, etc. A Bit Error calculation equation is provided to measure the accuracy of the extracted watermark.

Keywords : audio watermarking, empirical mode decomposition, hilbert transform, wavelet transform, adaptive threshold, binary image.

1. INTRODUCTION

In any given application, practitioners had already recognized two different technologies for multimedia data protection: encryption and digital watermarking [1, 2]. However, a digital watermark is complementary approach than cryptographic processes. The availability of digital multimedia contents has been grown rapidly in the recent years. Today, digital media documents can be distributed via the World Wide Web to a large number of people without much effort and money. Unlike traditional analog copying, with which the quality of the duplicated content is degraded, digital tools can easily produce large amount of perfect copies of digital documents in a short period.

This ease of digital multimedia distribution over the Internet, together with the possibility of unlimited duplication of this data, threatens the intellectual property (IP) rights more than ever. Therefore, there is a strong need for security services in order to keep ownership for the document owner and reliable to the customer. Watermarking plays an important role for that purposes. In a digital watermarking process, typically, a pattern or sequence of bits is embedded into a digital image, an audio or video file such that the embedded pattern can be detected or extracted later to make an assertion about the object. Some proposed or actual watermarking applications are [3]: broadcast monitoring, owner identification, proof of ownership, transaction tracking, content authentication and tampering detection, copy control, and device control, fingerprinting, information carrier.

Embedding secret information into audio sequences is a more tedious task than image watermarking, due to dynamic preeminence of the human auditory system (HAS) over human visual system (HVS). The main confront in digital audio watermarking and steganography is that if the perceptual transparency parameter is predetermined, the design of a watermark system cannot obtain high robustness and a high watermark data rate at the same time. We have addressed this challenging situation with the use of EMD specially developed for analyzing non-linear and non-stationary signals [4].

In the wavelet transform domain [5, 6], the watermark is embedded in high frequency band or in low frequency band; therefore watermark can't be embedded in the whole transform domain. In EMD-based method, we can add the watermark into the whole transform domain. The EMD method provides perfect time-frequency localization properties due to its instantaneous frequency than discrete wavelet transform (DWT) domain and leads to implicit audio-visual masking [7]. This decomposition method is adaptive, and highly efficient for perceptual transparency (inaudibility) and robustness. That is HAS is perfectly met too. In the previous EMD-based method [8], the watermark is added into the whole coefficients of the highest energetic IMF, therefore the watermarked IMF is sensitive to common signal processing attacks than this proposed method. This paper presents a method, which

Author α ¥ : Lecturer, Dept. of Computer Science & Engineering, Dhaka International University, Dhaka, Bangladesh.

E-mails : ibrahim_5403@yahoo.com, pramanikanis@gmail.com

Author σ : B.Sc Engineering from Dhaka International University major in Computer Science & Engineering. E-mail : abedinwasif@gmail.com

Author ρ ω : Lecturer, Dept. of Electrical, Electronics & Telecommunication Engineering, Dhaka International University, Dhaka, Bangladesh. E-mails : mahfuza.khatun@yahoo.com, tareq.mhd@gmail.com

can embed the watermark into the filtered coefficients of the highest energetic IMF using adaptive threshold. Since adaptive threshold is used to add watermark into the significant IMF, our method is much more resistant to common signal processing manipulation when compared with other methods.

II. PROPOSED AUDIO WATERMARKING TECHNIQUE

The Empirical Mode Decomposition is a new method for analyzing nonlinear and non-stationary data. This method is adaptive, and, therefore, highly efficient. Since this decomposition method is based on the local characteristic time scale of the data, it is applicable to nonlinear and non-stationary processes. Using this method any complicated data set can be decomposed into a finite and often small number of 'Intrinsic Mode Functions (IMF)'. An Intrinsic mode Function (IMF) is a function that satisfies two conditions [9]: (i) in the whole data set, the number of extrema and the number of zero crossings must either equal or differ at most by one; and (ii) at any point, the mean value of the envelope defined by the local maxima and the envelope defined by the local minima is zero. With the Hilbert transform, the 'Intrinsic Mode Functions' yield instantaneous frequencies as functions of time that provide sharp identifications of imbedded structures.

We have extended the previous idea by using EMD and embedding watermark into the filtered coefficients of the highest energetic IMF for increasing robustness against different signal processing attacks. In the proposed method, at first we have divided the host signal into a number of frames. Then each frame is decomposed into a small number of intrinsic mode functions (IMFs) that acknowledge well-behaved Hilbert transforms. This method is applicable to nonlinear and non-stationary process because of the local characteristic time scale of the data. However, after calculating energy of the every IMF, we have picked an IMF contained highest energy. The following Fig. 1 shows the energy distribution of the IMFs for the first frame of an audio signal:

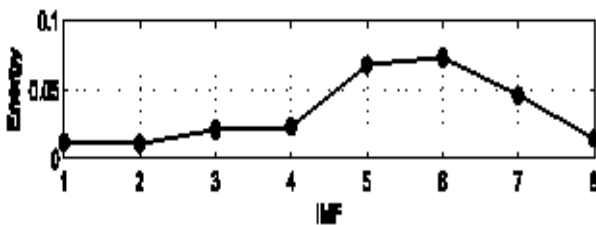


Figure 1 : Energy distribution for IMFs of the first frame

From the energy distribution of IMFs, we can select the 6th IMF because of the highest energetic value. A scaled version of the pseudo random

watermark sequence $w[n]$, n is the size of the watermark which is equal to the number of frame, added to the filtered coefficients of the selected IMFs. Only one bit of $w[n]$ is added to the only one frame. The i^{th} bit of $w[n]$ is added to the selected IMF (e.g. 6th IMF) of the i^{th} frame by using the following additive embedding formula (1):

$$d'_m(l) = d_m^i(l) + \alpha w(i) \quad (1)$$

Where, $i=1$ for the first frame, $i=2$ for the second frame and so on. l runs over all coefficients of the $\text{IMF} > T$ (Adaptive threshold), $d_m^i(l)$ is the filtered coefficient of m^{th} IMF of i^{th} frame, $w(i)$ is the i^{th} bit of the watermark, and $d'_m(l)$ is the watermarked coefficient of the selected IMF. The watermarked frame is generated by simply summing all IMF of that frame. The watermarked signal $x'[t]$ is generated after embedding watermarks to all frames in this way. The scaling factor α is chosen to be as high as possible while remaining inaudible. Here α is taken as 0.02. The threshold value (T) is calculated by this equation (2):

$$T = (\text{Max_Val} + \text{Min_Val}) / 2 \quad (2)$$

Where Max_Val is the maximum coefficient value and Min_Val is the minimum coefficient value of the selected IMF.

Since, EMD is empirical, there is no reason the decomposition will result in same (or even similar) IMFs when the signal is modified (or even the watermark added). If it does not result in the same IMFs, the entire scheme fails. Therefore, in the detection process, the original signal $x[t]$ and watermarked signal $x'[t]$ are transformed into Hilbert transform domain by excluding the use EMD. Then the watermark is extracted simply by using the following subtraction equation (3):

$$w'_i[t] = H'_i[x'[t]] - H_i[x[t]] / \alpha \quad (3)$$

Where, $H'_i[x'[t]]$ is the watermarked audio signal in Hilbert transform domain, $H_i[x[t]]$ is the original audio signal of i^{th} frame in Hilbert transform domain, and $w'_i[t]$ is the extracted watermark of the i^{th} frame. The resulted watermark bit is calculated by averaging over the number of embedded watermark to the frame. In this calculation, two different techniques are applied to two different watermarks.

For the sequence of 1 and -1, the watermark bit is a +1 If the average is greater than zero and if the

average is less than zero the corresponding watermark bit is a -1 . If we use binary number as a watermark, then the watermark bit is a 1 if the average is greater than zero; if the average is less than equal zero the corresponding watermark bit is a zero. Finally we get the resulted watermark $w''[n]$ after calculating the above operation for all frames.

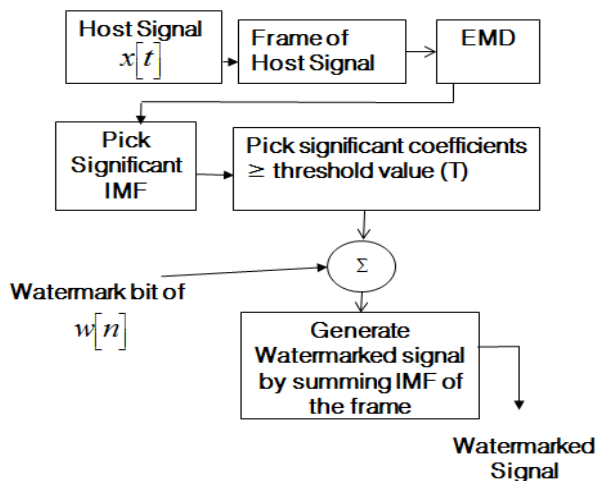
The similarity between the original watermark and resulted watermark is calculated to detect the existence of watermark. The similarity should be increased or decreased according to the amount of the inserted watermark, as well as to the degree of the attack. To calculate the detector response value (S_{dr}), dr means detector response, the proposed method utilizes the vector projection as defined in the following equation (4):

$$S_{dr} = w[n] \cdot w''[n] / \sqrt{w''[n] \cdot w''[n]} \quad (4)$$

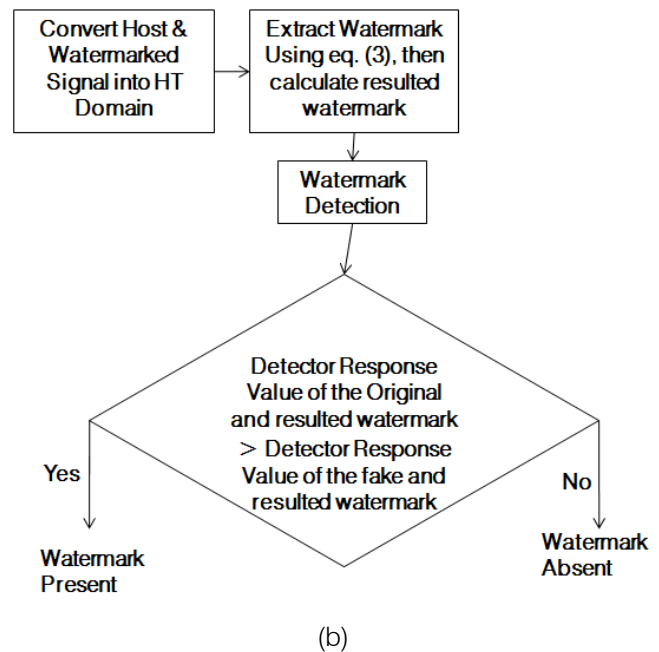
Where, $w[n]$ is the original watermark or fake watermark, and $w''[n]$ is the resulted watermark. We also calculated the normalized correlation (ψ) to measure the performance of the embedding algorithm using the following equation (5):

$$\psi = \left[\sum_j w(j)w''(j) / \sum_j w(j)w(j) \right] \quad (5)$$

When the detector response value of the original watermark and extracted watermark is greater than other detector response values of the fake watermarks and resulted watermark, then we can conclude that the watermark is present. Fig. 2 shows a block diagram of the proposed method for a single frame.



(a)

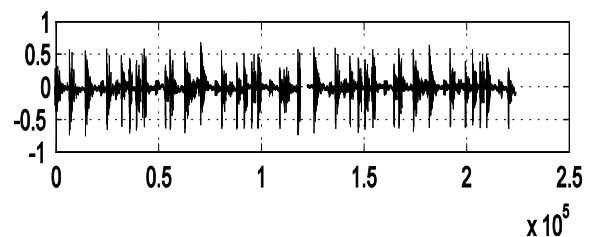


(b)

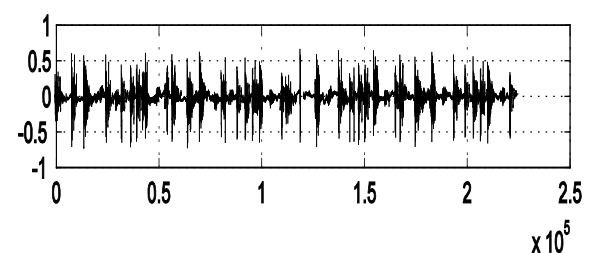
Figure 2 : The proposed method: (a) Watermark embedding and (b) Watermark detection

III. EXPERIMENTAL RESULTS AND DISCUSSION

To test the algorithm we have used eight audio clips (44 kHz, 705 kpbs), which has been collected from Fruity Loops software. An original audio signal (drum music), watermarked audio signal, and the extracted watermark signal are shown in Fig. 3. After converting the original and watermarked audio signal into HT domain, the extracted watermark is calculated using equation (3).



(a)



(b)

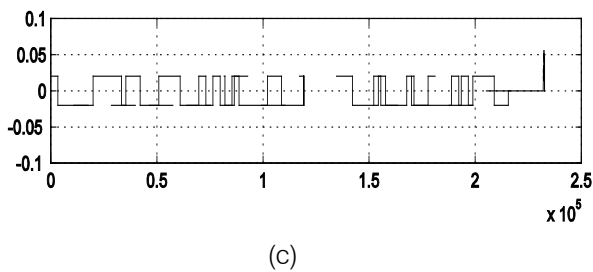


Figure 3 : (a) Original audio signal, (b) Watermarked audio signal, (c) Extracted Watermark

The detector response values (S_{dr}) on the y-axis and randomly generated watermark on the x-axis against three attacks as example for the piano music and drum music are shown in the following figures (from Fig. 4 to Fig. 6). The original audio signal is watermarked with a seed of 100. Therefore, in each figure, the detector response values with the original watermark are located at 100 on the x-axis. Therefore, at this position the detector response values are greater than any other values.

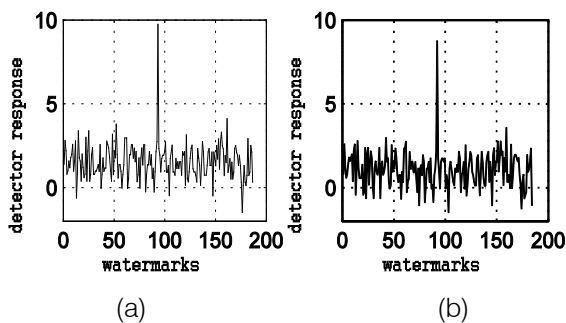


Figure 4 : Detector responses for Gaussian noise attack.

(a) Piano music with $\psi = 0.96$, SNR= -7.26 db and
(b) Drum music with $\psi = 0.9$, SNR= -15.104db

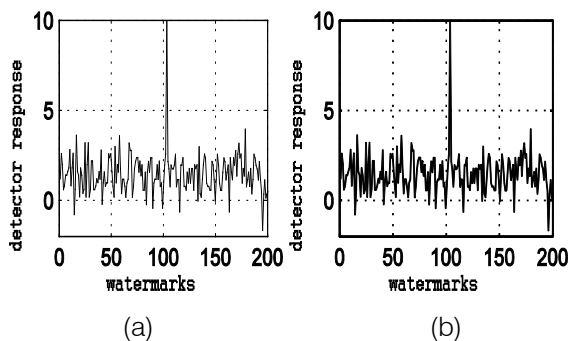


Figure 5 : Detector responses for Flanging effect

(a) Piano music with $\psi = 1$, SNR= 6.02db and (b) Drum music with $\psi = 1$, SNR= 6.05db

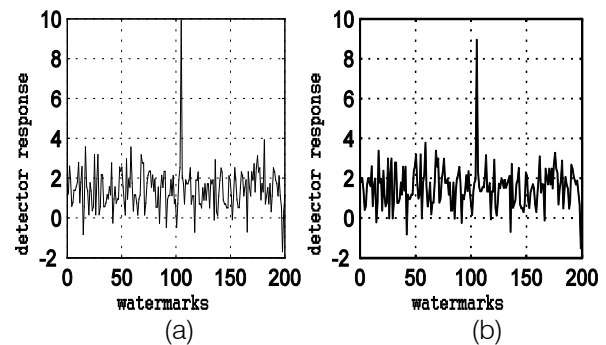


Figure 6 : Detector responses for MP3 compression
(a) Piano music with $\psi = 1$ and (b) Drum music with $\psi = 0.9$, 56 kbps

The performance against Gaussian noise addition is calculated to illustrate the robustness of the proposed algorithm as:

$$R_p = V_w - V_F \quad (6)$$

Where, V_w is the similarity measurement value between original watermark and resulted watermark extracted from the attacked watermarked signal and V_F is the highest similarity measurement value between false watermark and resulted watermark extracted from the attacked watermarked signal. The changes of the robustness performance for drum music are shown in Fig. 7:

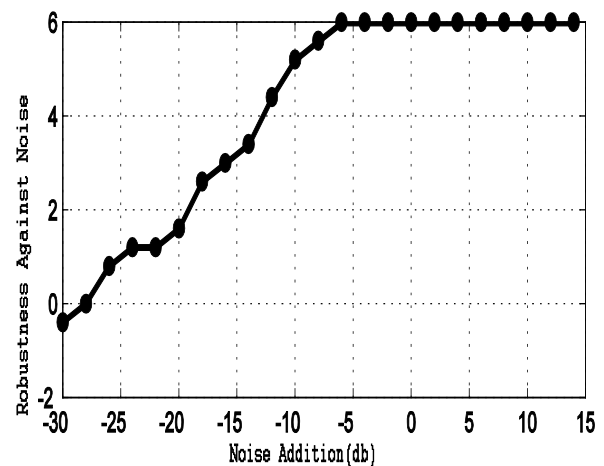


Figure 7 : Robustness performance values against different noise levels (db)

If $R_p > 0$, then watermark is detected correctly. The highest value of R_p represents maximum performance. The proposed method works for -26db also. We have also calculated Bit Error (β) to examine extracted watermark by the following equation (7):

$$\beta = \frac{\sum_j |w(j) - w''(j)|}{2L} \quad (7)$$

Where, L is the length of the watermark i.e. the number of bit of the watermark. The minimum value of

β is 0 and maximum value of β is 1. Experimental results of the eight music clips are shown in Table-I.

Table 1 : Experimental results of different attacks for various types of audio signals

Music Clip →		Drum	C_HC	Dream-zofluxury	NewStuff	ISuck-AtTitles	404lament	H441-GOA	Average β
Attacks ↓									
Gaussian noise	$\psi = 0.96$ SNR=-7.26 $\beta = 0.020$	$\psi = 0.90$ SNR=-15.10 $\beta = 0.050$	$\psi = 0.980$ SNR=-20.10 $\beta = 0.010$	$\psi = 0.94$ SNR=-4.23 $\beta = 0.030$	$\psi = 0.96$ SNR=-5.075 $\beta = 0.020$	$\psi = 0.94$ SNR=-7.36 $\beta = 0.030$	$\psi = 0.94$ SNR=-23.60 $\beta = 0.030$	$\psi = 0.94$ SNR=-17.25 $\beta = 0.030$	0.0275
Add FFT Noise	$\psi = 1$ SNR=30.5 $\beta = 0$	$\psi = 1$ SNR=28.5 $\beta = 0$	$\psi = 1$ SNR=30.2 $\beta = 0$	$\psi = 1$ SNR=35.5 $\beta = 0$	$\psi = 1$ SNR=40.5 $\beta = 0$	$\psi = 1$ SNR=64.21 $\beta = 0$	$\psi = 1$ SNR=48.14 $\beta = 0$	$\psi = 1$ SNR=54.46 $\beta = 0$	0.0
Down Sampling	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	$\psi = 1$ 22 kHz $\beta = 0$	0.0
Requanzization	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	$\psi = 1$ 352 kbps $\beta = 0$	0.0
Low pass filter	$\psi = 1$ SNR=28.68 $\beta = 0$	$\psi = 1$ SNR=16.68 $\beta = 0$	$\psi = 1$ SNR=17.62 $\beta = 0$	$\psi = 1$ SNR=31.88 $\beta = 0$	$\psi = 1$ SNR=20.81 $\beta = 0$	$\psi = 1$ SNR=22.19 $\beta = 0$	$\psi = 1$ SNR=24.83 $\beta = 0$	$\psi = 1$ SNR=30.39 $\beta = 0$	0.0
Delay Effect	$\psi = 1$ SNR=6.04 $\beta = 0$	$\psi = 1$ SNR=6.104 $\beta = 0$	$\psi = 1$ SNR=6.0309 $\beta = 0$	$\psi = 0.86$ SNR=6.116 $\beta = 0.070$	$\psi = 0.94$ SNR=6.094 $\beta = 0.030$	$\psi = 0.92$ SNR=6.16 $\beta = 0.040$	$\psi = 1$ SNR=6.12 $\beta = 0$	$\psi = 1$ SNR=6.12 $\beta = 0$	0.0175
Flanging Effect	$\psi = 1$ SNR=6.02 $\beta = 0$	$\psi = 1$ SNR=6.05 $\beta = 0$	$\psi = 1$ SNR=6.0167 $\beta = 0$	$\psi = 0.98$ SNR=6.185 $\beta = 0.01$	$\psi = 0.96$ SNR=6.051 $\beta = 0.020$	$\psi = 0.98$ SNR=6.096 $\beta = 0.010$	$\psi = 1$ SNR=6.061 $\beta = 0$	$\psi = 1$ SNR=6.05 $\beta = 0$	0.005
AllPass Reverberation	$\psi = 1$ SNR=5.77 $\beta = 0$	$\psi = 1$ SNR=3.88 $\beta = 0$	$\psi = 0.96$ SNR=4.0 $\beta = 0.020$	$\psi = 0.90$ SNR=5.495 $\beta = 0.050$	$\psi = 1$ SNR=4.343 $\beta = 0$	$\psi = 0.94$ SNR=3.905 $\beta = 0.030$	$\psi = 1$ SNR=6.76 $\beta = 0$	$\psi = 1$ SNR=2.80 $\beta = 0$	0.0125

The average bit errors of the eight music clips for different Gaussian noise addition are shown in Fig. 8:

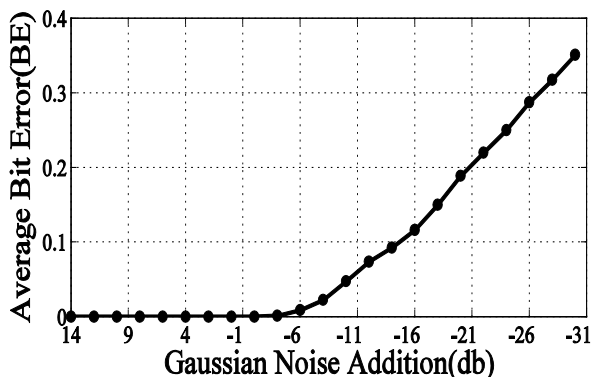


Figure 8 : Average bit-error against different level of Gaussian noise (db)

IV. CONCLUSION

This paper proposed an audio watermarking schema using the principle advantages of Empirical Mode Decomposition. A multiband approach is presented here for digital audio watermarking. The data adaptive time varying filtering is implemented for sub band decomposition of the host audio signal using EMD. Traditional frequency domain filtering techniques (i.e. Fourier and wavelet) depend on the priori basis functions and hence the signal is fitted with those bases. In the proposed method, the decomposition is fully data dependent and without the use of any basis function. The EMD based decomposition is a complete decomposition i.e. the original signal can easily be obtained by simply summing up the individual IMFs and

hence the method performs better for audio watermarking. The efficiency of the proposed algorithm has been shown by a series of experiment. In this method we can embed a binary logo image as a watermark also.

REFERENCES RÉFÉRENCES REFERENCIAS

1. M. Eskicioglu & E. J. Delp (April 2001). Overview of Multimedia Content Protection in Consumer Electronics Devices. *Signal processing: Image Communication*, 16(7), pp: 681-699.
2. A. M. Eskicioglu & E.J. Delp (April 2003). Security of Digital Entertainment Content From Creation to Consumption. *Signal Processing : Image Communication, Special Issue on Image Security*, 118 (4), pp. 237-262.
3. I. J. Cox, M. L. Miller, and J. A. Bloom (2002). Digital Watermarking. *Morgan Kaufmann Publishers*.
4. Huang, N. E. et. al. (1998). The empirical mode decomposition and the Hilbert spectrum for nonlinear and non-stationary time series analysis. *Proc. Roy. Soc. London A*, Vol. 454, pp. 903-995.
5. Xianghong Tang, Yamei Niu, Hengli Yue, Zhongke Yin (2005). A Digital Audio Watermark Embedding Algorithm. *International Journal of Information Technology, China*, Vol. 11 No.12.
6. Martinez-Noriega, R. et. al. (July 8-11, 2007). Audio Watermarking Based on Wavelet Transform and Quantization Index Modulation. *The 22nd International Technical Conference on Circuits/-Systems, Computers and Communications, ITC-CSCC 2007*, pp. 133-134, Busan, Korea.
7. LingFei Liang, ZiLiang Ping (2008). Information Hiding Based on Empirical Mode Decomposition. *2008 IEEE Pacific-Asia Workshop on Computational Intelligence and Industrial Application*, paciia, vol. 1, pp. 561-565.
8. A. N. K. Zaman, K. M. I. Khalilullah, M. W. Islam and M. K. I. Molla (May, 2010). A robust digital audio watermarking algorithm using empirical mode decomposition. *Proc. of 23rd Canadian Conf. on Electrical and Computer Engineering, (CCECE'10)*.
9. Foo S, Yeo T & Huang D. (2001). An adaptive audio watermarking system. *In: Proc. IEEE Region 10 International Conference on Electrical and Electronic Technology*, Phuket Island-Langkawi Island, Singapore, p 509-513.



Human Skin Detection

By Md. Zargis Talukder, Avizit Basak & Dr. Muhammad Shoyaib

Bangladesh University of Engineering and Technology

Abstract - Skin-color modeling is a crucial task for several applications of computer vision. Problems such as face detection in video are more likely to be solved if an efficient skin-color model is constructed. Most potential applications of skin-color model require robustness to significant variations in races, differing lighting conditions, textures and other factors. Given the fact that a skin surface reflects the light in a different way as compared to other surfaces. As the color of human skin is created by the combination of blood (red) and melanin (brown, yellow) which gives it a restricted range of hues. A skin region can be classified by comparing large image content of skin database and non-skin database. The RGB color space is widely used and most effective to detect skin region from an image. The segmentation is used to localize and identify homogeneous regions in a picture by perceptual attributes which include the size, the shape and the texture and/or color information. The probability of each RGB color space of skin and non-skin database is important to detect skin pixels.

Keywords : skin detection, image, processing.

GJCST-F Classification: 1.5.4



Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

Human Skin Detection

Md. Zargis Talukder^α, Avizit Basak^σ & Dr. Muhammad Shoyaib^ρ

Abstract - Skin-color modeling is a crucial task for several applications of computer vision. Problems such as face detection in video are more likely to be solved if an efficient skin-color model is constructed. Most potential applications of skin-color model require robustness to significant variations in races, differing lighting conditions, textures and other factors. Given the fact that a skin surface reflects the light in a different way as compared to other surfaces. As the color of human skin is created by the combination of blood (red) and melanin (brown, yellow) which gives it a restricted range of hues. A skin region can be classified by comparing large image content of skin database and non-skin database. The RGB color space is widely used and most effective to detect skin region from an image. The segmentation is used to localize and identify homogeneous regions in a picture by perceptual attributes which include the size, the shape and the texture and/or color information. The probability of each RGB color space of skin and non-skin database is important to detect skin pixels. For each pixel of testing image, the RGB value is calculated then the probability ratio of that RGB color space for training Skin and non-skin data base is compare with a threshold variable called θ . The threshold value lies between 0 and 1 but for this analyzing purpose threshold value has been taken as 0.4. If the ratio will be greater than 0.4 then that pixel will be detected as skin pixel else detected as non-skin pixel. After segmenting out skin from images, this can be useful for identifying faces, hand sign recognition, offensive content such pornography. The performance curve (ROC curve) reflects the overall accuracy of our analysis.

Keywords : skin detection, image, processing.

1. INTRODUCTION

Skin detection consists in detecting human skin pixels from an image. The system output is a binary image defined on the same pixel grid as the input image. Skin detection plays an important role in various applications such as face detection, searching and filtering image content on the web. Research has been performed on the detection of human skin pixels in color images by use of various statistical color models. Some researchers have used skin color models such as Gaussian, Gaussian mixture or histograms. In most experiments, skin pixels are

acquired from a limited number of people under a limited range of lighting conditions. Unfortunately, the illumination conditions are often unknown in an arbitrary image, so the variation in skin colors is much less constrained in practice. This is particularly true for web images captured under a wide variety of conditions. A good skin classifier must be able to discriminate between skin and non-skin pixels for a wide range of people with different skin types such as white, yellow, brown and dark and be able to perform well under different illumination conditions such as indoor, outdoor and with white and non-white illumination sources. An important step in the image classification process is color segmentation of the image into homogeneous skin-color regions and non-skin color regions in color space that is relatively invariant to minor illuminant changes.

However, given a large collection of labeled training pixels including all human skin (Caucasians, Africans, Asians) we can still model the distribution of skin and non-skin colors in the color space. Jones and Rehg [3] proposed techniques for skin color detection by estimating the distribution of skin and non-skin color in the color space using labeled training data. The histogram models were found to be slightly superior to Gaussian mixture models in terms of skin pixel classification for this color space. A skin detection system is never perfect and different users use different criteria for evaluation. General appearance of the skin-zones detected, or other global criteria might be important for further processing. For quantitative evaluation, false positives and detection rates have used. False positive rate is the proportion of non-skin pixels classified as skin and detection rate is the proportion of skin pixels classified as skin. The user might wish to combine these two indicators his own way depending on the kind of error he is more willing to afford. The larger the value ensures the larger the belief for a skin pixel. Error rates for all possible thresholding are summarized in the Receiver Operating Characteristic (ROC) curve.

In this research Image dataset is classified as Training Image Dataset and Testing Image dataset. The training Image dataset consist two database called Skin database and non-skin database consisting human skin pixels and non-skin pixels respectively. Whereas the testing image dataset containing real image both skin and non-skin pixels. The training skin database consists of five hundred and fifty five images and so as mask corresponds to the actual image. Each of them is

Author α : Department of Industrial and Production Engineering (IPE), Bangladesh University of Engineering and Technology (BUET), Bangladesh. E-mail : zargis_jarj@yahoo.com

Author σ : B.Sc in Electrical & Electronics Engineering from Rajshahi University of Engineering & Technology (RUET), Rajshahi, Bangladesh. E-mail : dhrubo_eee88@yahoo.com

Author ρ : Associate Professor, Institute of Information Technology (IIT), University of Dhaka (DU), Dhaka, Bangladesh. E-mail : shoyaib@univdhaka.edu

manually segmented such that the skin pixels are labeled carefully. The training non-skin database consists of Eight thousand four hundred and thirty images having no skin pixels. On the other hand testing image dataset contain One Hundred and three image. We select our image data set with proper illuminating condition as far possible. The goal is to infer a model from this set of data in order to perform skin detection on new images.

II. IMAGE PREPARATION

Most of the research efforts on skin detection have focused on visible spectrum imaging. Skin-color detection in visible spectrum can be a very challenging task as the skin color in an image is sensitive to various factors such as:

a) Illumination

A change in the light source distribution and in the illumination level (indoor, outdoor, highlights, shadows, non-white lights) produces a change in the color of the skin in the image (color constancy problem). The illumination variation is the most important problem among current skin detection systems that seriously degrades the performance.

b) Camera Characteristics

Even under the same illumination, the skin-color distribution for the same person differs from one camera to another depending on the camera sensor characteristics. The color reproduced by a CCD camera is dependent on the spectral reflectance, the prevailing illumination conditions and the camera sensor sensitivities.

c) Ethnicity

Skin color also varies from person to person belonging to different ethnic groups and from persons across different regions. For example, the skin color of people belonging to Asian, African, Caucasian and Hispanic groups is different from one another and ranges from white, yellow to dark.

d) Individual Characteristics

Individual characteristics such as age, sex and body parts also affect the skin-color appearance.

e) Other Factors

Different factors such as subject appearances (makeup, hairstyle and glasses), background colors, shadows and motion also influence skin-color appearance.

Many of the problems encountered in visual spectrum can be overcome by using non-visual spectrum such as infrared (IR) and spectral imaging. Skin-color in non-visual spectrum methods is invariant to changes in illumination conditions, ethnicity, shadows and makeup. However, the expensive equipment necessary for these methods combined with tedious

setup procedures have limited their use to specific application areas such as biomedical applications [1].

III. SKIN-COLOR MODELING AND ANALYSIS

a) Color Spaces

The choice of color space can be considered as the primary step in skin-color classification. The RGB color space is the default color space for most available image formats. Any other color space can be obtained from a linear or non-linear transformation from RGB. The color space transformation is assumed to decrease the overlap between skin and non-skin pixels thereby aiding skin-pixel classification and to provide robust parameters against varying illumination conditions.

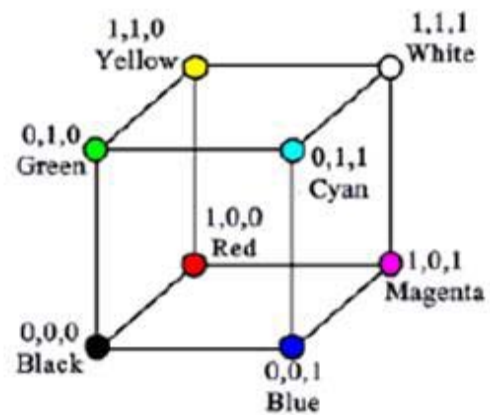


Figure 1 : RGB Color Space

The RGB color space consists of the three additive primaries: red, green and blue. Spectral components of these colors combine additively to produce a resultant color. The RGB model is represented by a 3-dimensional cube with red green and blue at the corners on each axis (Figure 1). Black is at the origin. White is at the opposite end of the cube. The gray scale follows the line from black to white. In a 24-bit color graphics system with 8 bits per color channel, red is (255, 0, 0). On the color cube, it is (1, 0, 0). The red, green and blue color components are highly correlated. The RGB model simplifies the design of computer graphics systems but is not ideal for all applications.

b) Explicit skin-color space thresholding

Image segmentation is a process of pixel classification into classes. Two classes or bi-level segmentation is the simplest case. The real world consists of many objects and so the number of classes. There are primarily four types of segmentation techniques: thresholding, boundary-based, region-based, and hybrid techniques. All these four methods can be implemented in grayscale and color images.

Thresholding is based on histogram of the image and applying assumption that peaks in the

histogram correspond to either background or objects of interest that can be extracted by separating histogram valleys. In segmentation, an image is partitioned into different homogeneous regions, where the homogeneity of a region may be composed based on different criteria such as gray level, color or texture. Region of interest segmentation plays a crucial role as a preliminary step for high-level image processing such as image analysis and pattern recognition.

Jones and Rehg [3], built a 3D RGB histogram model with two billion pixels collected from 18,696 web images. They reported that 77% of the possible RGB colors are not encountered and most of the histogram is empty. There is a marked skew in the distribution towards the red corner of the color cube due to the presence of skin in web images, though only 10% of the total pixels are skin pixels.

This suggested that skin colors occur more frequently than other object colors. By computing two different histograms, skin and non-skin histograms, the probability that a given color belongs to skin and non-skin class (also called class conditional probabilities) is defined as

$$P(\text{rgb} | \text{skin}) = \frac{S[\text{rgb}]}{T_s}$$

$$P(\text{rgb} | \text{non-skin}) = \frac{N[\text{rgb}]}{T_n}$$

Where $s[\text{rgb}]$ is the pixel count contained in bin "rgb" of the skin histogram, $n[\text{rgb}]$ is the equivalent count from the non-skin histogram, and T_s and T_n are the total counts contained in the skin and non-skin histograms, respectively. From the generic skin and non-skin histograms a reasonable separation can be obtained between skin and non-skin classes. This fact can be used to build fast and accurate skin classifiers even for images collected from unconstrained imaging environments such as web images, given that the training dataset is sufficiently huge. A larger training set can lead to better probability density function estimations. Given the class conditional probabilities of skin and non-skin-color models, a skin classifier can be built by which a given image pixel can be classified as skin, if

$$\frac{P(c/\text{skin})}{P(c/\text{non-skin})} \geq \theta$$

Where $0 \leq \theta \leq 1$ is a threshold value which can be adjusted to trade-off between true positives and false positives. This threshold value is normally determined from the ROC (receiver operating characteristics) curve calculated from the training data set. The ROC curve represents the relationship between the true positives and false positives as function of the detection threshold.

IV. DATA PREPARATION AND ANALYSIS

The color of skin in images depends primarily on the concentration of hemoglobin and melanin and on the conditions of illumination. It is well-known that the hue of skin is roughly invariant across different ethnic groups after the illuminant has been discounted. This is because differences in the concentration of pigments primarily affect the saturation of skin color, not the hue. For detecting skin a set of 9,088 real images are gathered. These images set have divided into 2 set of data called Training data set and testing data set. The training data have further divided into two dataset called Skin database and Non-skin database. The Skin database consists of 555 real images containing some skin pixels. Also 555 corresponding masked images are contained within this database. These images contain skin pixels belonging to persons of different origins, with unconstrained illumination and background conditions, which make the skin detection task more challenging and difficult. Figure 2 shows some images and their corresponding masked images. The skin pixels within these images have labeled manually and prepared for skin histogram. The masks are prepared manually with proper caring and consciousness. The Skin database sample consists of 20796952 image pixels and will be validated by a similar number of image pixels randomly selected. Each image pixel has been selected only once.

The Non-skin database consists of 8430 images that did not contain any skin pixel have prepared for the non-skin histogram. These images does not contain any skin pixels with unconstrained illumination and background conditions, The Non-Skin database sample consists of 800927510 image pixels for analyzing purpose. Each data file will be used to train a network during a training run.

The test data consists of 103 images selected at randomly from the remaining image database. The test images are used to evaluate the performance of the skin detection system.

For analyzing skin and non-skin dataset various programming language are available such as "C", "C++", "JAVA", "MATLAB" etc. All has flexibility enough and so as constrains. As JAVA is a platform independent language and has various built-in library function, so we choice JAVA for our analyzing purpose. Without it JAVA is one of the most powerful programming language this is another reason to select JAVA for our analyzing purpose.

a) Skin database analysis

The skin database has contained 555 real images containing skin pixel within them, which we considered as training image. There also have a mask database containing mask corresponding to those actual images. The mask was prepared manually.

During preparing mask the non-skin pixels were marked as white and skin as others. Hence by comparing actual image and corresponding to its mask the actual skin pixels have been detected easily. At the starting of the program we declared a three dimensional array having $256 \times 256 \times 256$ space for skin database to construct

a 256 bin histogram. When the skin pixel was detected, corresponding RGB color space have been increased by 1. The total process was continued till the last skin pixel detection in the Skin-database. After that the probability has calculated to each RGB color space which stored in a file named skin probability.



Figure 2 : Training Skin Database Image Corresponding to its Mask

b) Non-Skin Database Analysis

The Non-skin database has contained 8430 real image which does not contain any skin pixels. At the starting of the program we also declared a three dimensional array having $256 \times 256 \times 256$ space for Non-skin database to construct a 256 bin histogram. For each non skin pixels, the basic RGB combination is

detected. Then the array is incremented by 1 to its corresponding array space (RGB combination). The total process continued till to find out the last pixel in the Non-Skin database. After that the probability has calculated for each RGB space which stored in a file named non-skin probability.

The calculated general Information about Skin and Non-skin image are listed below:

	Total Counts Total	Total Occupied Bins	Total Unoccupied Bins	Percent Unoccupied
Skin Model	20796952	433460	16343748	97.4%
Non-skin Model	800927510	3237834	13539382	80.7%
Overlapping skin/non-skin bins	433394			
Skin pixels as a percentage of total pixels:	2.53 %			
Total photos in labeled dataset:	9088			
Percentage of photos containing skin:	7.15 %			

Table 1 : General analysis about skin and non-skin database

c) Histogram-based Skin Classifier

In this stage some binary images have been generated from its actual images. It is the testing stage, so testing dataset have used in this stage. Testing dataset contains 103 real images. The masks correspond to its real image of training database have prepared carefully. In this stage the first goal was to generate binary image from testing image using probability of RGB color space of training skin and non-skin databases, which we generate earlier. An important terminology called "threshold" is magnificent which is denoted by $\emptyset$ which can vary from 0 to 1. We compared

the value of $\emptyset$ with the ratio of skin probability and Non-skin Probability. So it is important consideration to fix up the value of $\emptyset$. As the variation of illumination condition, skin color depending on geographical location and many other unavoidable constrain threshold $\emptyset$ can be vary. So normally it depends on our image dataset. The general relationship between threshold $\emptyset$ and the ratio of skin and Non-skin probability can be defined as:

$$\frac{P(c/\text{skin})}{P(c/\text{non-skin})} \geq \emptyset$$

Where $0 < \theta < 1$.considering our training dataset we have taken 0.4 as the value of Threshold. From testing dataset we have considered each pixel and its relevant RGB values. Depending on that RGB value the skin probability value and Non-skin probability value have chosen from calculated dataset. If the ratio of skin probability to the non-skin probability corresponds to the testing RGB space is greater than 0.4 then that

pixel considered as a skin pixel and so that the output marker defined as 1 that is Pure white, otherwise we considered as a non-skin pixel and so that the output marker is defined as 0 that is pure Black. By this process a binary image corresponds to its actual training images is generated. In figure 3 some testing image, mask and generated binary images are shown.

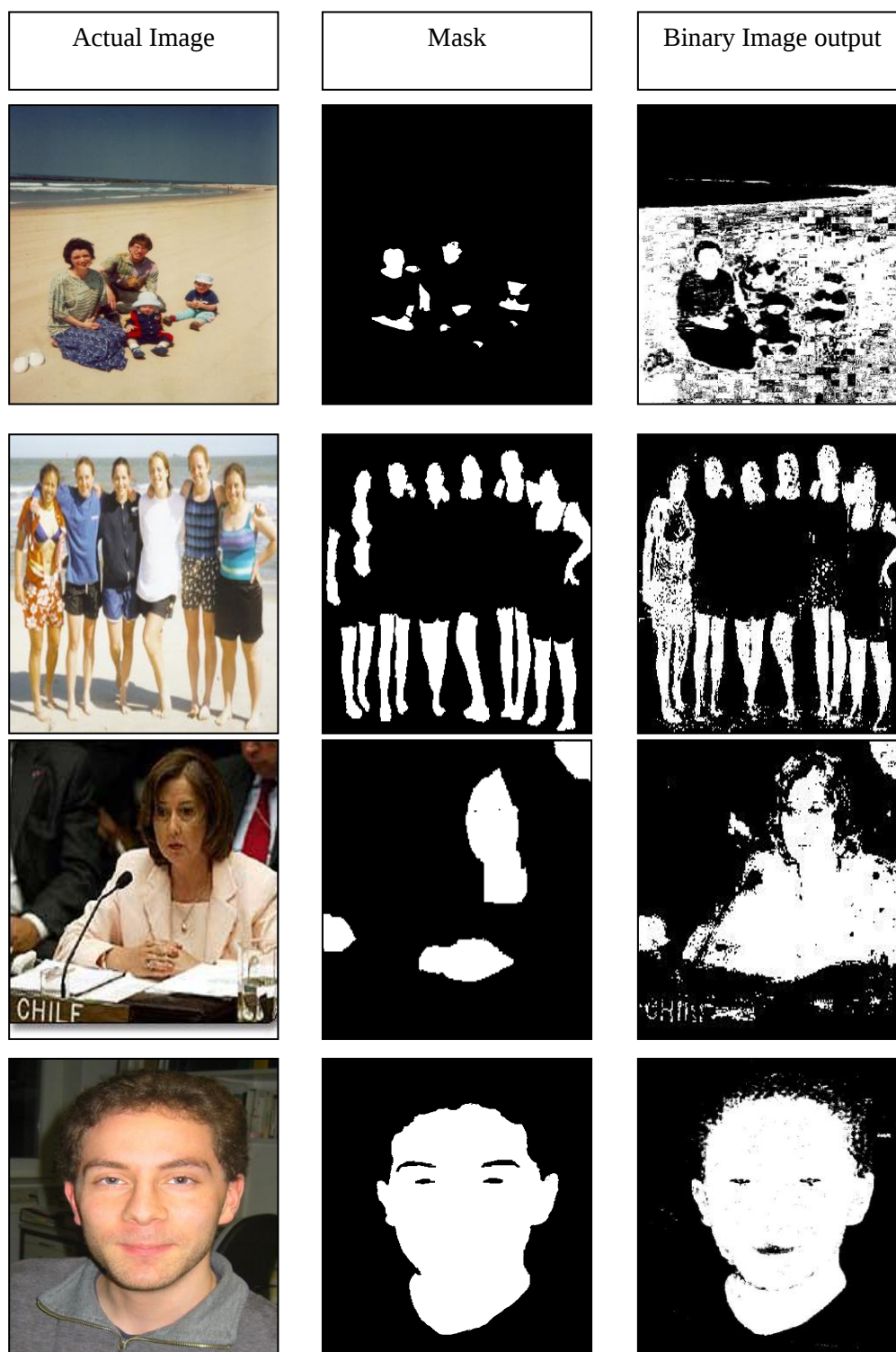


Figure 3 : Testing image, mask and generated binary image

The classifier does a good job of detecting skin in most images. In particular, the skin labels form dense sets whose shape often resembles that of the true skin pixels. The detector tends to fail on highly saturated or shadowed skin. The example photos also show the performance of the detector on non-skin pixels. In photos such as the sand or flowers (which having Human skin like color) the false detections are sparse and scattered. More problematic are images with wood or copper-colored metal such as the kitchen scene or railroad tracks. These photos contain colors which often occur in the skin model and are difficult to discriminate reliably. This results in fairly dense sets of false positives. Classifier performance can be quantified by computing the ROC curve which measures the threshold-dependent trade-off between misses and false detections.

V. PERFORMANCE ANALYSIS

The ROC (Receiver Operating Characteristic) curve is an important trend-off for measuring our performance. In addition to the threshold setting, classifier performance is also a function of the size of the histogram (number of bins) in the color models. Too few bins results in poor accuracy while too many bins lead to over fitting. Figure 4 shows the family of ROC curves produced as the size of the histogram varies from 256 bins/channel. The axis labeled "Probability of correct detection" gives the fraction of pixels labeled as skin that were classified correctly, while "Probability of false detection" gives the fraction of non-skin pixels which are mistakenly classified as skin. These curves were computed from the test data.

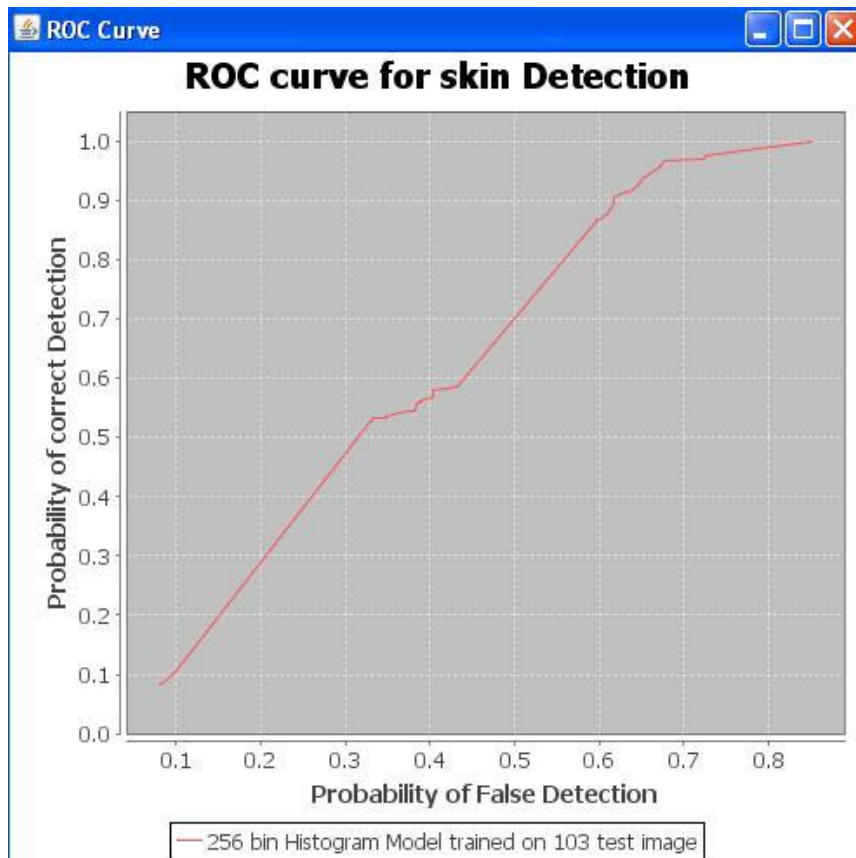


Figure 4 : ROC Curve for 256 Bin Histogram Model

By analyzing other size of bin, it can be found that histogram size 32 gave the best performance, superior to the size 256 model at the larger false detection rates and slightly better than the size 16 model in two places. The performance of the skin classifier is good considering the unconstrained nature of web images. The classifier (256 bin) can detect roughly $(4258578 / (4258578 + 3470391)) * 100\% = 55\%$ of skin pixels with a false positive rate of 45%. This corresponds to the point on the ROC curve where the

probability of false rejection (which is one minus the probability of correct detection) equals the probability of false detection.

In addition to histogram size, classifier performance is also affected by the amount of training data. To do so the list of skin and non-skin images in the training set and divided then into chunks containing approximately 5744138.0 skin pixels and 27375162 non-skin pixels. On each iteration we added one such chunk of new skin and non-skin pixels to the evolving training

set. A ROC curve was computed at each iteration showing the classifier performance on the full test set.

If more data is added, performance on the training set decreases because the overlap between skin and non-skin data increases. In that case a smooth shaped curve can be obtained which is a measure of good performance and tends to fill up the area near to 1 underlining the curve. Increasing testing data increase the performance characteristic in a similar manner.

VI. CONCLUSION

Skin detection is very important but very complex procedures are maintained. A single model is can't give the proper result. But various mixture models can give more appropriate result. Whatever this analysis emphasis on analyze the basic properties of skin and fix up a proper threshold. The testing data set was small so that ROC curve occupied less area in the graph. Using 256 bins Histogram model is another reason to decrement the curve sharpness. Generally 32 bin color model can give more appropriate result. Many objects in the real world have skin-tone colors, such as some kinds of leather, sand, wood, fur, etc., which might be mistakenly detected by a skin detector. Therefore, skin detection can be very useful in finding human faces and hands in controlled environments where the background is guaranteed not to contain skin-tone colors. Since skin detection depends on locating skin-colored pixels, its use is limited to color images, i.e., it is not useful with gray-scale, infrared, or other types of image modalities that do not contain color information. This research hypothesize that taking a minimalistic approach to the software design that can perform image preprocessing, and skin detection in real time.

REFERENCES RÉFÉRENCES REFERENCIAS

1. P. Kakumanu, S. Makrogiannis and N. Bourbakis. A survey of skin-color modeling and detection methods. *ELSEVIER, Science Direct, Pattern Recognition* 40, pages 1106 – 1122, (2007).
2. Sanjay Kr. Singh, D. S. Chauhan, Mayank Vatsa, Richa Singh. A Robust Skin Color Based Face Detection Algorithm. *Tamkang Journal of Science and Engineering*, Vol. 6, No. 4, pages 227-234 (2003).
3. Michael J. Jones., James M. Rehg. *Statistical Color Models with Application to Skin Detection*, December 1998.
4. E. Ahmed., C. Muang and D. Hu., *Skin Detection*, Rutgers University, Piscataway, NJ, 08902, USA.
5. M.H. Yang, D.J. Kriegman, N. Ahuja, *Detecting faces in images: a survey*, *IEEE Trans. Pattern Anal. Mach. Intell.* 24 (1) pages 34 – 58. (2002).
6. W. Zhao, R. Chellappa, P. J. Philips, A. Rosenfeld, *Face recognition: a literature survey*, *ACM Comput. Surveys* 85 (4) pages 299 – 458. (2003).

7. S.G. Kong, J. Heo, B.R. Abidi, J. Paik, M.A. Abidi, *Recent advances in visual and infrared face recognition - a review*, *Comput. Vision Image Understanding* 97 (2005).
8. C. Balas, *An imaging colorimeter for noncontact tissue color mapping*, *IEEE Trans. Biomed. Eng.* 44 (6) pages 468 - 474. (1997).
9. Richard O. Duda and Peter E. Hart. *Pattern Classification and Scene Analysis*. John Wiley and Sons, 1973.
10. Brian Funt, Kobus Barnard, and Lindsay Martin. *Is machine color constancy good enough?* In Hans Burkhardt and Bernd Neumann, editors, *Proceedings of 5th European Conference on Computer Vision*, volume I, pages 445–459, Freiburg, Germany, June 1998.



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

Accessed · 24, 25, 34
Adaptive · 38, 40, 42, 43, 44, 46
Almighty · 23
Anova · 27, 28, 29, 31, 32, 34, 35, 37
Appropriate · 28, 31, 32, 35, 56

B

Binary · 10, 13, 19, 21, 38, 42, 46, 48, 52, 54

C

Catastrophic · 15
Classification · 17, 19, 21, 23, 24, 25, 26, 57
Classifier · 17, 19, 21, 23, 24, 25, 26, 52, 55
Compression · 10, 12, 14, 15
Configuration · 5
Consecutive · 12
Containing · 48, 50, 51, 52, 55
Converting · 10, 42
Counterpart · 27, 28, 29, 31, 32, 34, 35, 37

D

Decomposition · 38, 39, 40, 44, 46
Detection · 25, 48, 49, 50, 52, 54, 55, 56

E

Efficiency · 14, 17, 46
Empirical · 38, 40, 46
Establishment · 5, 6
Extension · 27, 31, 34, 35

F

Frequency · 10, 12, 13, 38, 39, 44

G

Graphical · 27, 28, 29, 31, 32, 34, 35, 37
Grayscale · 49
Guaranteed · 56

H

Histogram · 5, 6, 48, 49, 50, 51, 52, 55
Huffman · 10

I

Including · 17, 19, 21, 23, 25, 26
Invariant · 19, 48, 49, 50

J

Jittered · 27, 28, 29, 31, 32, 34, 35, 37

M

Measurement · 3, 43
Microelectronic · 1

N

Nonlinearly · 21

O

Obtained · 38, 44, 49, 50, 56

P

Parallel · 28
Parametrical · 8

R

Researcher · 35
Robust · 15, 38, 40, 42, 43, 44, 46, 56

S

Sensors · 1, 2, 3, 5, 6, 8
Simulation · 1, 2, 3, 5, 6, 8

Statistical · 12, 14, 27, 32, 48
Subtraction · 40

T

Threshold · 24, 25, 38, 40, 42, 43, 44, 46, 54
Toolbox · 1, 2, 3, 5, 6, 8
Trajectories · 25

U

Unavoidable · 53

V

Vector · 17, 19, 21, 23, 24, 25, 26
Visualize · 27, 29, 35
Voltages · 2, 3

W

Watermarking · 10, 12, 14, 15, 38, 40, 42, 43, 44, 46
Weeds · 17, 19, 21, 23, 25, 26
Workshop · 25, 46

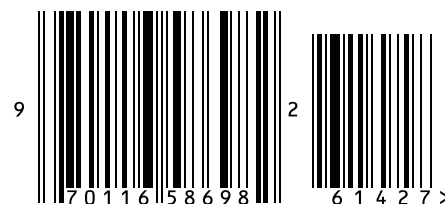


save our planet



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