

GLOBAL JOURNAL

OF COMPUTER SCIENCE AND TECHNOLOGY

DISCOVERING THOUGHTS AND INVENTING FUTURE

Technology
Reforming
Ideas

December 2011

Pinnacles

Hidden Markov Model

Implementing 3D Warping

Segmentation of Calculi

Augmented Reality Method

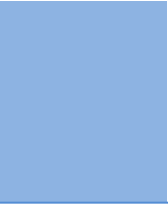
The Volume 11

Issue 22

VERSION 1.0



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY

VOLUME 11 ISSUE 22 (VER. 1.0)

OPEN ASSOCIATION OF RESEARCH SOCIETY

© Global Journal of Computer Science and Technology.2011.

All rights reserved.

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

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

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

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

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

Engage with the contents herein at your own risk.

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

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

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

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

Global Journals Inc.

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

Sponsors: Global Association of Research

Open Scientific Standards

Publisher's Headquarters office

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

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

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

Offset Typesetting

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

Packaging & Continental Dispatching

Global Journals, India

Find a correspondence nodal officer near you

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

eContacts

Press Inquiries: press@globaljournals.org

Investor Inquiries: investors@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

Luis Galárraga

J!Research Project Leader

Saarbrücken, Germany

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

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. Automatic Raaga Identification System for Carnatic Music Using Hidden Markov Model. ***1-9***
 2. Implementing 3D Warping Method In Wavelet Domain. ***11-15***
 3. Representing Aspect Model as Graph Transformation. ***17-24***
 4. Extended Apriori for association rule mining: Diminution based utility weightage measuring approach. ***25-30***
 5. Data mining with Predictive analysis for healthcare sector: An Improved weighted associative classification approach. ***31-36***
 6. Fuzzy and Swarm Intelligence for Software Cost Estimation. ***37-41***
 7. Segmentation of Calculi from Ultrasound Kidney Images by Region Indicator with Contour Segmentation Method. ***43-51***
 8. An Efficient and Speedy Activity Model for Information System Based Organizations. ***53-58***
 9. A Computer Vision Based Collaborative Augmented Reality Method for Human-Computer Interaction. ***59-63***
 10. Headlight Prefetching for Cooperative Media Streaming in Mobile Environments. ***65-70***
 11. Implementation of Impulse Noise Reduction Method to Color Images using Fuzzy Logic. ***71-75***
 12. A Classification of Aerial Data Based on Data Mining Clustering Algorithm. ***77-81***
-
- vii. Auxiliary Memberships
 - viii. Process of Submission of Research Paper
 - ix. Preferred Author Guidelines
 - x. Index



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Automatic Raaga Identification System for Carnatic Music Using Hidden Markov Model

By Prasad Reddy P.V.G.D, B. Tarakeswara Rao, Dr. K.R Sudha, Hari CH.V.M.K.

Andhra University, Visakhapatnam, India

Abstract - As far as the Human Computer Interactions (HCI) is concerned, there is a broad range of applications in the area of research in respect of Automatic Melakarta Raaga Identification in music. The pattern of identification is the main object for which, the basic mathematical tool is utilized. On verification, it is observed that no model is proved consistently and effectively to be predicted in its classification. This paper is, therefore, introduces a procedure for Raaga Identification with the help of Hidden Markov Models (HMM) which is rather an appropriate approach in identifying Melakarta Raagas. This proposed approach is based on the standard speech recognition technology by using Hidden continuous Markov Model. Data is collected from the existing data base for training and testing of the method with due design process relating to Melakarta Raagas. Similarly, to solve the problem of automatic identification of raagas, a suitable approach from the existing database is presented. The system, particularly, this model is based on a Hidden Markov Model enhanced with Pakad string matching algorithm. The entire system is built on top of an automatic note transcriber. At the end, detailed elucidations of the experiments are given. It clearly indicates the effectiveness and applicability of this method with its intrinsic value and significance.

Keywords : Melakarta Raaga, Raaga Recognition, Hidden Markov Models (HMM) classifier, Pakad.

GJCST Classification : G.3, I.2.6



AUTOMATIC RAAGA IDENTIFICATION SYSTEM FOR CARNATIC MUSIC USING HIDDEN MARKOV MODEL

Strictly as per the compliance and regulations of:



Automatic Raaga Identification System for Carnatic Music Using Hidden Markov Model

Prasad Reddy P.V.G.D.^α, B. Tarakeswara Rao^Ω, Dr. K.R Sudha^β, Hari CH.V.M.K.^ψ

Abstract - As for as the Human Computer Interactions (HCI) is concerned, there is broad range of applications in the area of research in respective of Automatic Melakarta Raaga Identification in music. The pattern of identification is the main object for which, the basic mathematical tool is utilized. On verification, it is observed that no model is proved consistently and effectively to be predicted in its classification.

This paper is, therefore, introduces a procedure for Raaga Identification with the help of Hidden Markov Models (HMM) which is rather an appropriate approach in identifying Melakarta Raagas. This proposed approach is based on the standard speech recognition technology by using Hidden continuous Markov Model. Data is collected from the existing data base for training and testing of the method with due design process relating to Melakarta Raagas. Similarly, to solve the problem of automatic identification of raagas, a suitable approach from the existing database is presented.

The system, particularly, this model is based on a Hidden Markov Model enhanced with Pakad string matching algorithm. The entire system is built on top of an automatic note transcriber. At the end, detailed elucidations of the experiments are given. It clearly indicates the effectiveness and applicability of this method with its intrinsic value and significance.

Keywords : Melakarta Raaga, Raaga Recognition, Hidden Markov Models (HMM) classifier, Pakad.

I. INTRODUCTION AND PROBLEM DEFINITION

Raaga classification has been a central pre-occupation of Indian music theorists for centuries [1], reflecting the importance of the concept in the musical system. The ability to recognize raagas is an essential skill for musicians and listeners. The raaga provides a framework for improvisation, and thus, generates specific melodic expectations for listeners, crucially impacting their experience of the music. Recently, MIR researchers have attempted to create systems that can accurately classify short raaga excerpts [1, 2]

Raaga recognition is a difficult problem for humans, and it takes years for listeners to acquire these skills for a large corpus. It can be very difficult to precisely explain precisely the essential qualities of a raaga. Most common descriptions of raagas are not sufficient to distinguish them. For example, many raagas

share the same notes, and even similar characteristic phrases. While ascending and descending scales are sometimes used absolutely in an artistic manner. They are highly abstracted from the real melodic sequences found in performances. It takes a performer lengthy practice to fully internalize the raaga and be able reproduces it. This is despite the fact that humans are highly adept at pattern recognition, they have little problem with pitch recognition, even in harmonically and timbrally dense settings. Clearly, there are several other difficulties for an automatic Melakarta raaga recognition system.

The special feature of the paper is the study of musical raagas through major contributions. In first place, our solutions based primarily on techniques from speech processing and pattern matching, which shows that techniques from other domains can be purposefully extended to solve problems in computational musical raagas. Secondly, the two note transcription methods presented are novel ways to extract notes from sample raagas of Indian classical music. This approach has given very encouraging results. Hence these methods could be extended to solve similar problems in musical raagas and other domains.

The remaining parts of the paper is organized as follows. Section 2 highlights some of the useful and related previous research work in the area. The solution strategy is discussed in detail in Section 3. The test procedures and experimental results are presented in Section 4... Finally, Section 5 lists out the conclusions.

II. LITERATURE REVIEW

On verification of different concepts it is noticed that a that very little work has taken place in the area of applying techniques from computational musicology and artificial intelligence to the realm of Indian classical music. Of special interest to us, is the work done by Gaurav Pandey et al. [17] present an approach to solve the problem of automatic identification of Raagas from audio samples is based on a Hidden Markov Model enhanced with a string matching algorithm, Sahasrabuddhe et al. [7,8,6,14]. In their work, Raagas have been modelled as finite automata which were constructed by using information codified in standard texts on classical music. This approach was used to generate new samples of the Raaga, which were technically correct and were indistinguishable from compositions made by humans.

Author^α : Department of CS&SE, Andhra University, Visakhapatnam, India.

*Author^Ω : School of Computing, Vignana University, Guntur, India.
E-mail : tarak7199@gmail.com*

Author^β : Department of EEE, Andhra University, Visakhapatnam, India.

Author^ψ : Department of IT, Gitam University, Visakhapatnam, India.

It is also further noticed that Hidden Markov Models [6,18,19] are now widely used to model signals whose functions however are not known. A Raaga too, can be considered to be a class of signals and can be modeled as an HMM. The advantage of this approach is the similarity and it has with the finite automata formalism suggested above [11].

It is obvious that "Pakad" is a catch-phrase of the Raaga, with each Raaga having a different Pakad[19]. Most people claim that they identify the Raaga being played by identifying the Pakad [19] of the Raaga[17]. However, it is not necessary for a Pakad be sung without any breaks in a Raaga performance. Since the Pakad is a very liberal part of the performance in itself, standard string matching algorithms were not guaranteed to work. Approximate string matching algorithms designed specifically for computer musicology, such as the one by Iliopoulos and Kurokawa [10] for musical melodic recognition with scope for gaps between independent pieces of music, seemed more relevant.

The other connected works which deserve mention here are the ones on Query by Humming [8] and [9], and Music Genre Classification [7]. Although these approaches are not followed, there is a lot of scope for using such low level primitives for Raaga identification, and this might open avenues for future research.

III. PRESENT WORK

To solve the problems in speech processing Hidden Markov Models have been traditionally used. One important class of such problems is that involving word recognition. The problem is very closely related to the word recognition problem. This correspondence can be established by the simple observation that raaga compositions.

This correspondence between the word recognition and raaga identification problems is exploited to devise a solution to the latter. This solution is explained below. Also presented is an enhancement to this solution using the Pakad of a raaga. However, both these solutions assume that a note transcriber is readily available to convert the input audio sample into the sequence of notes used in it. It is generally cited in literature that "Monophonic note transcription is a trivial problem". However, our observations in the field of Indian classical music were counter to this, particularly because of the permitted variability in the duration of use of a particular note. To handle this, two independent heuristic strategies are designed for note transcription from any given audio sample. These strategies are explained in the corresponding pages.

a) Hidden Markov Models

The doubly embedded stochastic process is Hidden Markov Models[17,19] where the underlying stochastic process is not directly observable, and have

the capability of effectively modelling statistical variations in spectral features i.e. processes which generate random sequences of outcomes according to certain probabilities. More concretely, an HMM[19] is a finite set of states, each of which is associated with a (generally multidimensional) probability distribution. Transitions among the states are governed by a set of probabilities called transition probabilities. In a particular state, an outcome or observation can be generated, according to the associated probability distribution. It is only the outcome not the state that is visible to an external observer. So states are "hidden" and hence the name hidden Markov model.

In order to define an HMM [16,17,18,19] completely, the following elements are needed [12].

- The number of states of the model, N
- The number of observation symbols in the alphabet, M.
- A set of state transition probabilities

$$A = \{a_{ij}\} \quad (1)$$

$$a_{ij} = p\{q_{t+1} = j | q_t = i\}, 1 \leq i, j \leq N \quad (2)$$

where q_t denotes the current state.

- A probability distribution in each of the states,

$$B = \{b_j(k)\} \quad (3)$$

$$b_j(k) \geq p\{\alpha_t = v_k | q_t = j\}, 1 \leq j \leq N, 1 \leq k \leq M \quad (4)$$

where v_k denotes the k th observation symbol in the alphabet and α_t the current parameter vector. The initial state distribution, $\pi = \{\pi_i\}$ where,

$$\pi_i = p\{q_1 = i\}, 1 \leq i \leq N \quad (5)$$

Thus, an HMM can be compactly represented as

$$\lambda = (A, B, \pi) \quad (6)$$

The HMM[19] and their derivatives have been widely applied to speech recognition and other pattern recognition problems [1]. Most of these applications have been inspired by the strength of HMMs, ie the possibility to derive understandable rules, with highly accurate predictive power for detecting instances of the system studied, from the generated models. This also makes HMMs the ideal method for solving Raaga identification problems, the details of which are presented in the next subsection.

b) HMM in Raaga Identification

The raaga identification problem falls largely in the set of speech processing problems. This justifies the use of Hidden Markov Models in our solution[19]. Two other important reasons which motivated the use of HMM in the present context are:

1. The sequences of notes for different Raagas are very well defined and a model based on discrete states with transitions between them is the ideal representation for these sequences [4].
2. The notes are small in number, hence making the setup of an HMM easier than other methods.
3. This HMM, which is used to capture the semantics of a Raaga, is the main component of our solution. Construction of the HMM Used The HMM used in our solution is significantly different from that used in, say, word recognition. This HMM, which is called λ from now on that can be specified by considering each of its elements separately.
4. Each note in each octave represents one state in λ . Thus, the number of states, $N=12 \times 3=36$. Here, the three octaves of Indian classical music are considered, namely the Mandra, Madhya and Tar Saptak, each of which consists of 12 notes.
5. The transition probability $A = \{a_{ij}\}$ represents the probability of note j appearing after note i in a note sequence of the Raaga represented by λ . The initial state probability $\pi = \{\pi_i\}$ represents the probability of note i being the first note in a note sequence of the raaga represented by λ .

The outcome probability $B = \{B_i(j)\}$ is set according to the following formula

$$B_i(j) = \begin{cases} 0, & i \neq j \\ 1, & i = j \end{cases} \quad (7)$$

It is Observed, at each state α in λ , the only possible outcome is note α the last condition takes the hidden character away from λ , but it can be argued that this setup success for the representation of Raagas, as our solution distinguishes between performances of distinct Raagas on the basis of the exact order of notes sung in them and not on the basis of the embellishments used. HMM is also a statistical approach, and hence requires large amount of training data for effective estimate of the model parameters.

The system performance degrades when training and testing environments differ.

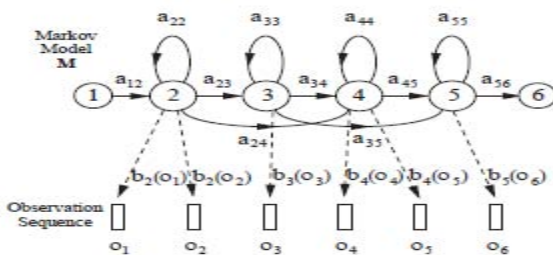


Fig. 1 : The Markov Generation Model

By applying the HMMs for Identification One such HMM λ_l , whose construction is described above, is set up for each Raaga l in the consideration set. Each of

these HMMs is trained, i.e. its parameters A and π (B has been pre-defined) is calculated using the note sequences available for the corresponding Raaga with the help of the Baum-Welch learning algorithm.

If all the HMMs have been trained, identifying the closest Raaga in the consideration set on which the input audio sample is based, is a trivial task. The sequence of notes representing the input is passed through each of the HMMs constructed, and the index of the required Raaga can be calculated as under:

$$\text{Index} = \text{argmax}(\log p(O|\lambda_l)), 1 \leq l \leq N \quad (8)$$

To complete the task, the required raaga can be determined as RaagaIndex. This preliminary solution gave reasonable results in our experiments (Refer to Section 4). However, there was still, a need to improve the performance by incorporating knowledge into the system. This can be done through the Pakad approach, which is discussed in the next section.

c) Pakad Matching

There is general expectation notion in the AI community that the incorporation of knowledge related to the problem being addressed into a system enhances its ability to solve the problem. One such very powerful piece of information about a Raaga is its Pakad[17,19].

Pakad is defined as a condensed version of the characteristic arrangement of notes, peculiar to each Raaga, when repeated in a recital from; it enables a listener to identify the Raaga being played. In other words, Pakad is a string of notes characteristic to a Raaga in which a musician frequently returns while improvising in a performance. The Pakad also serves as a springboard for improvisational ideas; each note in the Pakad can be embellished and improvised around to form new melodic lines. One common example of these embellishments is the splitting of Pakad into several substrings and playing each of them in order in disjoint portions of the composition, with the repetition of these substrings permitted. In spite of such permitted variations, Pakad is a major identifying characteristic of a Raaga and is used even by experts of Indian classical music for identifying the Raaga been played.

The characteristic elements of the Pakad give way to a string matching approach to solve the Raaga identification problem. Pakad matching can be used as reinforcement for an initial estimation of the underlying Raaga in a composition. Two ways of matching the Pakad with the input string of notes are devised in order to strengthen the estimation done. The incorporation of this step makes the final identification process a multi-stage one.

As indicated in the proceeding paragraphs δ -Occurrence with α -Bounded Gaps As mentioned earlier, the Pakad has to appear within the performance of a Raaga. However, it rarely appears in one segment as a whole. It is more common for it to be spread out, with substrings repeated and even other notes inserted in

between. This renders simple substring matching algorithms mostly insufficient for this problem. A more appropriate method for matching the Pakad is the δ -occurrence with α -bounded gaps algorithm. The algorithm employs dynamic programming and matches individual notes from the piece to be searched, say t , identifying a note in the complete sample, say p , as belonging to t only if:

1. There is a maximum difference of δ between the current note of p and the next note of t .
2. The position of occurrence of the next note of t in p is displaced from its ideal position by at most α . However, this algorithm assumes that a piece(t) can be declared present in a sample(p) only if all notes of the piece(t) are present in the sample(p) within the specified bounds. This may not be true in our case because of the inaccuracy of note transcription (Refer to 3.4). Hence, for each Raaga I in the consideration set, a score γI is maintained as

$$\gamma I = \frac{mI}{nI}, 1 \leq I \leq N_{\text{Raagas}} \quad (9)$$

mI = maximum number of notes of the Pakad of Raaga I identified nI = number of notes in the Pakad of Raaga I . This score is used in the final determination of the Raaga.

n-gram Matching Another method of capturing the appearance of the Pakad within a Raaga performance is to count the frequencies of appearance of successive n-grams of the Pakad. Successive n-grams of a string are its substrings starting from the beginning and going on till the end of the string is met. Also, to allow for minor gaps between successive notes, each n-gram is searched in a window of size $2n$ in the parent string. Based on this method, another score is maintained according to the formula,

$$\text{score}_I = \sum_n \sum_j \text{freq}_{j,n,I} \quad (10)$$

$\text{freq}_{j,n,I}$ = number of times the j^{th} n-gram of the Pakad of Raaga I is found in the input. This score is also used in the final determination of the underlying Raaga.

Final Determination of the Underlying Raaga

Once the above scores have been calculated, the final identification process is a three-step one.

1. The probability of likelihood probl is calculated for the input after passing it through each HMM λI and the values so obtained are sorted in increasing order. After reordering the indices as per the sorting, if

$$\frac{\text{Prob}_{N_{\text{Raaga}}} - \text{Prob}_{N_{\text{Raaga}-1}}}{\text{Prob}_{N_{\text{Raaga}-1}}} > n \quad \text{----- (11)}$$

then,

$$\text{Index} = N_{\text{Raagas}}$$

2. Otherwise, the values γI are sorted in increasing order and indices set accordingly. After this arrangement, if

$$\text{Prob}_{N_{\text{Raaga}}} > \text{Prob}_{N_{\text{Raaga}-1}}, \text{ and } \frac{Y_{N_{\text{Raaga}}} - Y_{N_{\text{Raaga}-1}}}{\text{Prob}_{N_{\text{Raaga}-1}}} > n$$

$$\text{then, Index} = N_{\text{Raagas}}$$

Otherwise, the final determination is made on the basis of the formula, where K is a predefined constant

$$\text{Index} = \text{argmax} (\log p(o|\lambda_I) + K * \text{score}_I), 1 \leq I \leq N \quad (12)$$

This procedure is used for the identification of the underlying Raaga. The steps enable the system to take into account all probable features for Raaga identification, and thus display good performance. The performance of the final version is discussed in Section 4.

d) Note Transcription

It is based on the assumption that the input audio sample has already been converted into a string of notes[17,19]. However, a few hurdles are faced in our efforts. The main hurdle in this conversion with regard to Indian classical music is the fact that notes are permitted to be spread over time for variable durations in any composition. Here, two heuristics are presented based on the pitch of the audio sample, which are used to derive notes from the input. They are very general and can be used for any similar purpose. From the pitch behavior of various audio clips, two important characteristics of the pitch structure are observed, based on the following two heuristics.

1. input sample.
2. The Hill Peak Heuristic[17,19].

This heuristic identifies notes in an input sample on the basis of hills and peaks occurring in the pitch graph. A sample pitch graph is shown in the figure 2.

A discreet study of an audio clip and its pitch graph shows that the notes occur at points in the graph where there is a complete reversal in the sign of the slope, and in many cases, also when there isn't a reversal in sign, but a significant change in the value of the slope. Translating this into mathematical terms, given a sample with time points $t_1, t_2, \dots, t_{i-1}, t_i, t_i+1, \dots, t_n$

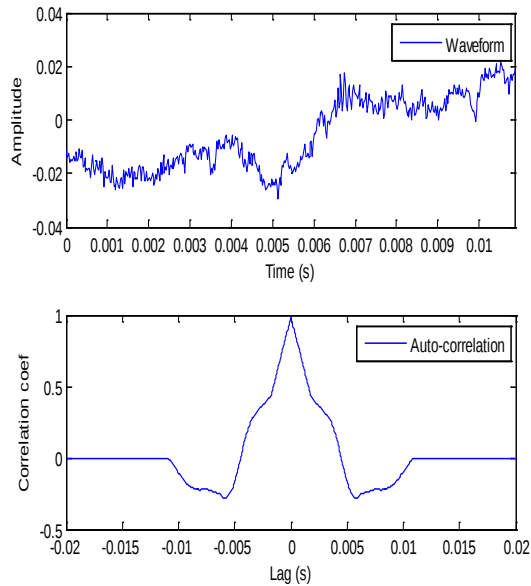


Fig2 : Simple Pitch Graph

Once the point of occurrence of a note has been determined, the note can easily be identified by finding the note with the closest characteristic pitch value. Performing this calculation over the entire duration of the sample gives the string of notes corresponding to it. An important point to note here is that, unless the change in slope between the two consecutive pairs of time points is significant, it is assumed that the last detected note is still in progress, thus allowing for variable durations of notes.

The Note Duration Heuristic is based on the assumption that in a composition a note continues for at least a certain constant span of time, which depends on the kind of music considered. Corresponding notes are calculated for all pitch values available. A history list of the last k notes identified is maintained including the current one (k is a pre-defined constant). The current note is accepted as a note of the sample, only if it is different from the dominant note in the history, i.e. the note which occurs more than m times in the history (m is also a constant). Sample values of k and m are 10 and 8. By making this test, a note can be allowed to extend beyond the time span represented by the history. A pass over the entire set of pitch values, gives the set of notes corresponding to the

IV. RESULTS AND DISCUSSION

In conclusion, there are four main modules in this paper. They are the extracted indicators of raaga features, training the features using HMM classifier, training HMM through multiple raaga data and testing the selected raagas by using the features of raaga. The results of classification obtained through both the features are combined to produce more accurate results. The regions commonly identified in both the classification results are now highlighted. In the

proposed model, the HMM classifier, Pakad matching, that was implemented to distinguish the various types of raagas is discussed. The implementation is done in MATLAB Speech Tollbox [20].

72 raagas are trained and also tested the said 72 raagas, all of them are recognized. Then, the raagas are tested which are not in the train. After testing the entire process, the obtained results are indicated in the given sample tables with the supporting evidences.

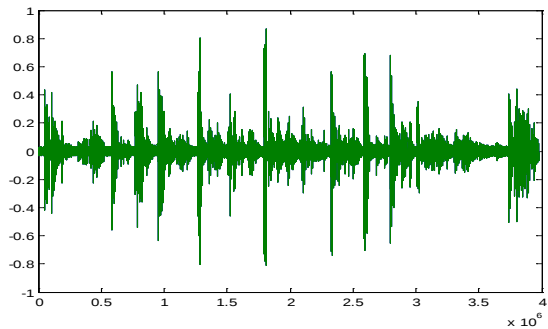


Fig 3 : Plot file for begada raaga (90 sec. Input Raaga)

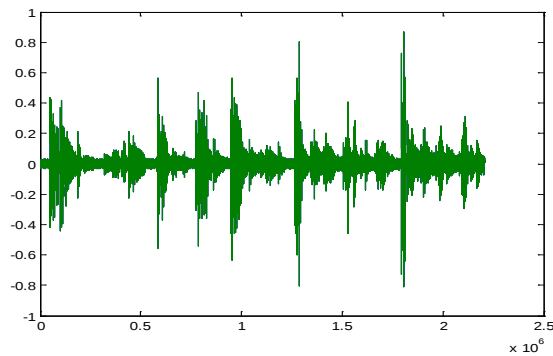


Fig 4 : Plot file for begada raaga (50 sec for Trained raaga)

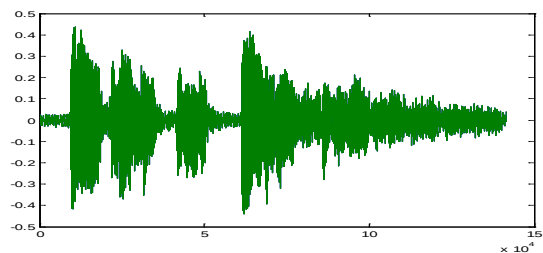


Fig 5 : Plot files for begada raaga (3sec for tested raaga)

Testing Time: First 3 seconds

Table 1 : Confusion matrix indicating that the first 3 sec of each raagas are tested.

Stimulus	Recognized Raagas (%)						
	B eg ad a	Van asa pat hi	sund avino dini	Desh	Kap inar aya ni	Ma ya mal ava go wla	Kan dana kuth uhal am
Begada	88	54	58	62	65	65	75
Vanasapa thi	54	88	63	70	72	75	76
sundavin odini	58	68	88	68	70	68	68
Desh	62	70	88	68	75	78	72
Kapinara yani	65	72	70	85	88	80	82
Mayamal avagowla	65	75	68	78	80	88	80
Kandana kuthuhalam	88	76	68	72	82	80	85

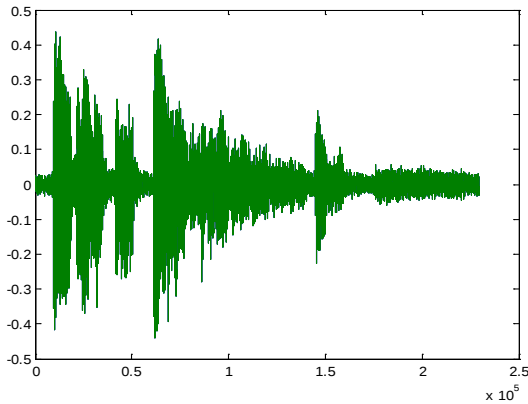


Fig 6 : Plot file for begada raaga (5sec for testing)

Testing Time: First 5 seconds

Table 2 : Confusion matrix indicating that the first 5 sec of each raagas are tested.

Stimulus	Recognized Raagas (%)						
	Be ga da	Va nas apa thi	sund avino dini	De sh	Kap inar aya ni	Ma yam alav ago wla	Kan dana kuth uhal am
Begada	90	58	58	62	65	65	75
Vanasap athi	58	90	63	70	72	75	76
sundavi nodini	58	68	90	68	70	68	68
Desh	62	70	90	68	75	78	72
Kapinara yani	65	72	70	85	90	80	82
Mayama lavagowla	65	75	68	78	80	90	80
Kandana kuthuhal am	90	76	68	72	82	80	88

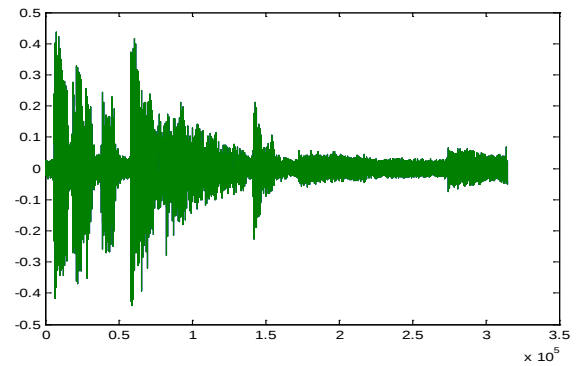


Fig 7 : Plot file for begada raaga (7sec for testing)

Testing Time : First 7 seconds

Table 3 : Confusion matrix indicating that the first 7 sec of each raagas are tested.

Stimulus	Recognized Raagas (%)						
	B eg ad a	V a n as a p at hi	sun dav ino dini	Des h	Ka pin ara yan i	M ay am ala va go wl a	Kan dan akut huh ala m
Begada	92	58	58	62	65	65	75
Vanasapa thi	58	92	63	70	72	75	76
sundavin odini	58	68	92	68	70	68	68
Desh	62	70	88	92	75	78	72
Kapinara yani	65	72	70	85	92	80	82
Mayamal avagowla	65	75	68	78	80	92	80
Kandana kuthuhal am	92	76	68	72	82	80	88

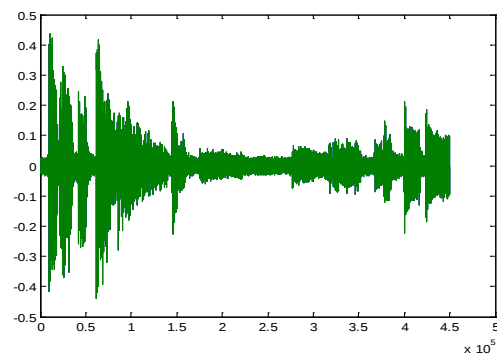


Fig 8 : Plot file for begada raaga (10sec for testing)

Testing Time : First 10 seconds

Table 4 : Confusion matrix indicating that the first 10 sec of each raagas are tested.

Stimu lus	Recognized Raagas (%)						
	Be ga da	Va na sa pa thi	sun davi nodi ni	Des h	Ka pin ara yani	May am ala va go wla	Kan dan akut huh ala m
Begada	92	58	58	62	65	65	75
Vanas apathi	58	92	63	70	72	75	76
sund avino dini	58	68	92	68	70	68	68
Desh	62	70	88	92	75	78	72
Kapin arayani	65	72	70	85	92	80	82
Maya malav agowla	65	75	68	78	80	92	80
Kand anakuth uhalam	92	76	68	72	82	80	88

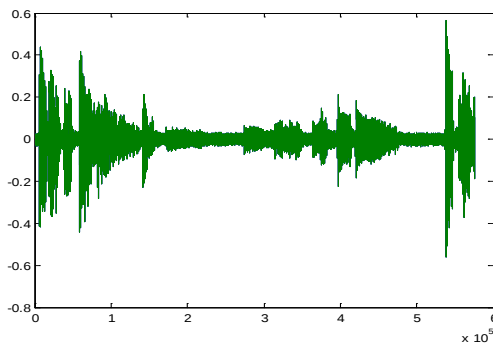


Fig 9 : Plot file for begada raaga (13sec for testing)

Testing Time: First 13 seconds:

Table 5 : Confusion matrix indicating that the first 13 sec of each raagas are tested.

Stimulus	Recognized Raagas (%)						
	Beg ada	V an as ap athi	sund avin odini	Des h	Kap inar aya ni	Ma ya ma lava gowla	Kan dana kuth uhal am
Begada	92	58	58	62	65	65	75
Vanasap athi	58	92	63	70	72	75	76
sundavi nodini	58	68	92	68	70	68	68
Desh	62	70	76	92	75	78	72
Kapinar ayani	65	72	70	85	92	80	82
Mayama lavagowla	65	75	68	78	80	92	80
Kandana kuthuhal alam	90	76	68	72	82	80	92

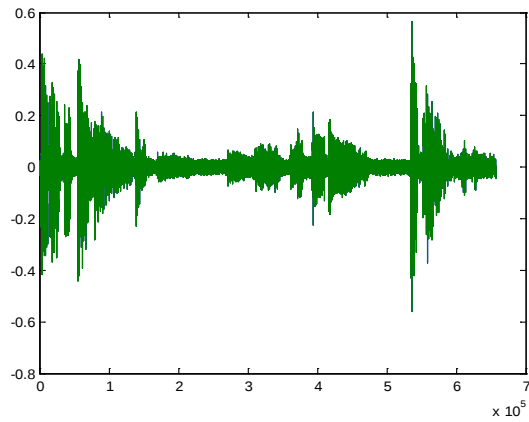


Fig 10 : Plot file for begada raaga (15sec for testing)

Testing Time: First 15 seconds

Table 6 : Confusion matrix indicating that the first 15 sec of each raagas are tested.

Stimu lus	Recognized Raagas (%)						
	Be ga da	Van asap athi	sund avino dini	Des h	Ka pin ara yan i	Ma ya ma lav agow la	Kan dan akut huh ala m
Begad a	92	58	58	62	65	65	75
Vanas apathi	58	92	63	70	72	75	76
sunda vinod ini	58	68	92	68	70	68	68
Desh	62	70	76	92	75	78	72
Kapin araya ni	65	72	70	85	92	80	82
Maya malav agowl a	65	75	68	78	80	92	80
Kanda nakut huh alam	90	76	68	72	82	80	92

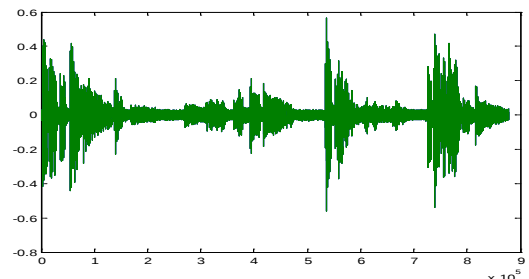


Fig 11 : Plot file for begada raaga (20sec for testing):

Testing Time: First 20 seconds

Table 7: Confusion matrix indicating that the first 20 sec of each raagas are tested.

Stimulus	Recognized Raagas (%)						
	Beg ada	Van asa path i	sund avin odin i	Des h	Kapi naray ani	Ma ya ma lav ago wla	Kan dan akut huh ala m
Begada	92	58	58	62	65	65	75
Vanasapat hi	58	92	63	70	72	75	76
sundavin odini	58	68	92	68	70	68	68
Desh	62	70	76	92	75	78	72
Kapinaray ani	65	72	70	85	92	80	82
Mayamal avagowla	65	75	68	78	80	92	80
Kandanak uthuhala m	90	76	68	72	82	80	92

V. CONCLUSION AND FUTURE WORK

The system for automatic raaga identification is presented in this paper it is based on Hidden Markov Models (HMM), Mel Frequency Cepstral Coefficients (MFCCs) and string matching. A very important part of the system is its note transcriptor, for which two (heuristic) strategies based on the pitch of sound are proposed. The strategy adopted is significantly different from those adopted for similar problems in Western and Indian classical music. Our problem, however, is different. Hence, probabilistic automata constructed on the basis of the notes of the composition are used to achieve our goal.

An attempt has also been made to determine the effectivity of Hidden Markov Model classifier and string matching algorithm is discussed for raaga recognition system that was developed. Raagas from instruments that are under four categories have been considered. The recognized raagas are presented in a confusion matrix table based on samples collected. As detailed in the experimental results, a high degree of accuracy of nearly 92% is achieved for recognized raagas of trained set where as an accuracy of around 70% reported for other raagas from outside sets. The HMM classifier, thus, achieved a better performance by in recognizing raagas from trained inputs, but able to differentiate when it the input parameters are other raagas from outside. although the approach used in the system is very general, there are two important directions for future research. In first place, a major part of the system is based on heuristics. There is a need to build this part on more rigorous theoretical foundations. Secondly, the constraints on the input for this system

are quite restrictive. The two most important problems which must be solved are estimation of the base frequency of an audio sample and multiphonic note identification. Solutions to these problems will help improve the performance and scope of the system.

REFERENCES

1. Bhatkande.V (1934), "Hindusthani Sangeet Paddhati.Sangeet Karyalaya", 1934.
2. Chordia.P (2006), Automatic raag classification of pitch-tracked performances using pitch-class and pitch-class dyad distributions", In Proceedings of International Computer Music Conference, 2006.
3. D. Weenink: "Praat: doing phonetics by computer": Institute of Phonetic Sciences, University of Amsterdam (www.praat.org)
4. M. Choudhary and P. R. Ray: "Measuring Similarities Across Musical Compositions: An Approach Based on the Raga Paradigm": Proc. International Workshop on Frontiers of Research in Speech and Music, pp. 25- 34.
5. L. R. Rabiner: "A Tutorial on Hidden Markov Models and Selected Applications in Speech Recognition": Proc. IEEE, Vol. 77, No. 2, pp. 257-286: February 1989.
6. C. S. Iliopoulos and M. Kurokawa: "String Matching with Gaps for Musical Melodic Recognition": Proc. Prague Stringology Conference, pp. 55-64: 2002.
7. H. V. Sahasrabudhe: "Searching for a Common Language of Ragas": Proc. Indian Music and Computers: Can 'Mindware' and Software Meet?: August 1994.
8. R. Upadhye and H. V. Sahasrabudhe: "On the Computational Model of Raag Music of India": Workshop on AI and Music: 10th European Conference on AI, Vienna: 1992.
9. G. Tzanetakis, G. Essl and P. Cook: "Automatic Musical Genre Classification of Audio Signals": Proc. International Symposium of Music Information Retrieval, pp. 205-210: October 2001.
10. Ghias, J. Logan, D. Chamberlin and B. C. Smith: "Query by Humming – Musical Information Retrieval in an Audio Database": Proc. ACM Multimedia, pp. 231-236: 1995.
11. H. Deshpande, U. Nam and R. Singh: "MUGEC: Automatic Music Genre Classification": Technical Report, Stanford University: June 2001.
12. S. Dixon: "Multiphonic Note Identification": Proc. 19th Australasian Computer Science Conference: Jan-Feb 2003.
13. W. Chai and B. Vercoe: "Folk Music Classification Using Hidden Markov Models": Proc. International Conference on Artificial Intelligence: June 2001.
14. J. Viterbi: "Error bounds for convolutional codes and an asymptotically optimal decoding algorithm": IEEE Transactions on Information Theory, Volume IT-13, pp.260-269: April 1967.

15. L. E. Baum: "An inequality and associated maximization technique in statistical estimation for probabilistic functions of Markov processes": Inequalities, Volume 3, pp. 1-8: 1972.
16. L. E. Baum and T. Petrie: "Statistical inference for probabilistic functions of finite state Markov chains": Ann.Math.Stat., Volume 37, pp. 1554-1563: 1966.
17. Gaurav Pandey, Chaitanya Mishra, and Paul Ipe "TANSEN: a system for automatic raga identification"
18. Prasad et al. "Gender Based Emotion Recognition System for Telugu Rural Dialects using Hidden Markov Models" Journal of Computing: An International Journal, Volume2 ,Number 6 June 2010 NY, USA, ISSN: 2151-9617
19. Tarakeswara Rao B et. All "A Novel Process for Melakartha Raaga Recognitionusing Hidden Markov Models (HMM)", International Journal of Research and Reviews in Computer Science (IJRRCS), Volume 2, Number 2, April 2011, ISSN: 2079-2557.
20. L.Arslan, Speech toolbox in MAT LAB, Bogazici University, <http://www.busim.ee.boun.edu.tr/~arslan/>



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Implementing 3D Warping Method In Wavelet Domain

By Indrit Enesi, Elma Zanj, Betim Çiço

Polytechnic University of Tirana, Albania

Abstract - A wide class of operations on images can be performed directly in the wavelet domain by operating on coefficients of the wavelet transforms of the images and other matrices defined by these operations. Operating in the wavelet domain enables one to perform these operations progressively in a coarse-to-fine fashion, operate on different resolutions, manipulate features at different scales, and localize the operation in both the spatial and the frequency domains. Performing such operations in the wavelet domain and then reconstructing the result is also often more efficient than performing the same operation in the standard direct fashion. Performing 3D warping in the wavelet domain is in many cases faster than their direct computation. In this paper we demonstrate our approach both on still and sequences of images.

Keywords : 3D warping, wavelet, multiresolution, planar, cylindrical, spherical, temporal coherence.

GJCST Classification : G.1.2



Strictly as per the compliance and regulations of:



Implementing 3D Warping Method In Wavelet Domain

Indrit Enesi^α, Elma Zanj^α, Betim Çiço^β

Abstract - A wide class of operations on images can be performed directly in the wavelet domain by operating on coefficients of the wavelet transforms of the images and other matrices defined by these operations. Operating in the wavelet domain enables one to perform these operations progressively in a coarse-to-fine fashion, operate on different resolutions, manipulate features at different scales, and localize the operation in both the spatial and the frequency domains. Performing such operations in the wavelet domain and then reconstructing the result is also often more efficient than performing the same operation in the standard direct fashion. Performing 3D warping in the wavelet domain is in many cases faster than their direct computation. In this paper we demonstrate our approach both on still and sequences of images.

Keywords : 3D warping, wavelet, multiresolution, planar, cylindrical, spherical, temporal coherence.

I. IMAGE-BASED RENDERING AND 3D WARPING

Many image-based rendering algorithms use pre-rendered or pre-acquired reference images of a 3D scene in order to synthesize novel views of the scene. The central computational component of such algorithms is 3D image warping, which performs the mapping of pixels in the reference images to their coordinates in the target image. In this paper we present wavelet warping — a new class of forward 3D warping algorithms for image based rendering. We rewrite the 3D warping equations as a point wise quotient of linear combinations of matrices. Rather than computing these linear combinations in a standard manner, we first pre-compute the wavelet transforms of the participating matrices. Next, we perform the linear combinations using only the unique non-zero wavelet transform coefficients. Applying the inverse wavelet transform to the resulting coefficients yields the desired linear combinations. We describe in detail wavelet warping algorithms for three common types of 3D image warps: planar-to-planar, cylindrical-to-planar, and spherical-to-planar. Current viewers allow the user to interactively change the viewing direction [1]. By using depth information, a 3D warper enables users to change the viewing position (center of projection), in addition to the viewing direction [2]. A fast 3D warper enables users to view a scene interactively. We will show that the wavelet

warping algorithm is at least as fast as the most efficient warping algorithm known to date for planar and cylindrical warps, and is nearly twice as fast in the spherical case. Perhaps more importantly, our wavelet warping algorithms support progressive multiresolution rendering. Considering an object whose image-based model consists of one or more high-resolution reference images, the high resolution may be necessary for a close-up view of the object, but for most views of a 3D scene containing the object a much lower resolution suffices. Our approach makes it possible to perform the warp at the appropriate coarser resolution, without unnecessarily warping every pixel in the reference images. Multi-resolution warping can also be achieved within a standard warping framework by using an over-complete pyramid-based image representation (e.g. a Laplacian pyramid), but at a cost of increasing the size of the representation [3]. Multiresolution wavelet warping has the advantage that the computation is progressive: a low resolution result can be progressively refined without redundant computations. We present a new algorithm for warping an entire sequence of images with depth to a novel view. This algorithm is also based on wavelet warping, and it utilizes the temporal coherence typically present in image sequences or panoramic movies to achieve considerable speedups over frame-by-frame warping.

II. CHOICE OF WAVELET TRANSFORM

There are two main requirements that a wavelet transform should satisfy in order to be suitable for our framework [4]:

1. The transform should be sparse.
2. Reconstruction (inverse wavelet transform) should be fast to compute.

To achieve faster reconstruction we choose transforms with smaller support size, and therefore fewer vanishing moments. This rules out the 9-7 transform which is considerably slower than the other transforms [5]. In particular, this transform requires floating point arithmetic, whereas the other transforms can be implemented using only integer additions and shifts. The S+P and TS transforms are similar. They are both special cases of the same transform, which is factored into the S transform followed by an additional lifting step, but with different prediction coefficients [6]. For our purposes it is sufficient to experiment with the

Author^{αβ} : Polytechnic University of Tirana, Albania.
E-mails : ienesi @fti.edu.al, ezanj@fti.edu.al, bcico@fti.edu.al

more efficient TS transform. In order to assess the speed and the sparsity of the remaining three wavelet bases (S, TS, and I(2, 2)) we gathered the relevant statistics over a database of 300 photography images representing landscapes, buildings, people, products, etc.. Each image was transformed from RGB to YIQ color space and processed at full (640x384) and at half (320x192) resolutions. The results are summarized in Tab. 1 and 2 and plots in Fig. 2. Our experiments indicate that all three transforms provide roughly the same sparsity of wavelet domain representation for natural images. We note that the percentage of

remaining coefficients is typically higher when operating on the half-resolution versions of the image. Decreasing the resolution results in smaller smooth regions in the images, and applying a transform with few vanishing moments yields fewer near-zero coefficients over these regions. In terms of speed, the S and I(2, 2) transforms are the fastest (the S transform is slightly faster), while the TS transform is slower by a factor of roughly 2 [7]. Consequently, the S transform was chosen for wavelet domain image blending and for wavelet domain convolution.

Tab. 1 : A comparison between the S, TS, and I(2,2) transforms For each transform and each image resolution we list the mean reconstruction time in milliseconds and the mean percentage of remaining (non-zero) coefficients, and the standard deviation corresponding to each mean.

Transform	S transform		TS transform		I(2,2) transform	
Resolution	640x384	320x192	640x384	32x192	640x384	320x192
Rec. Time	12.8 (0.7)	3.2 (0.3)	27.7 (1.4)	6.9 (0.7)	13.4 (0.7)	3.3 (0.3)
% non-zero	50.3 (13.7)	66.4 (12.4)	50 (12.8)	64.9 (12.3)	51.7 (11.2)	65.4 (10.9)

Tab. 2 : The average number of distinct non-zero values in a wavelet-transformed image for each of the Y, I,Q channels.

Transform/Channel	Y	I	Q
none	153	51	26
S	134	36	19
TS	121	32	18

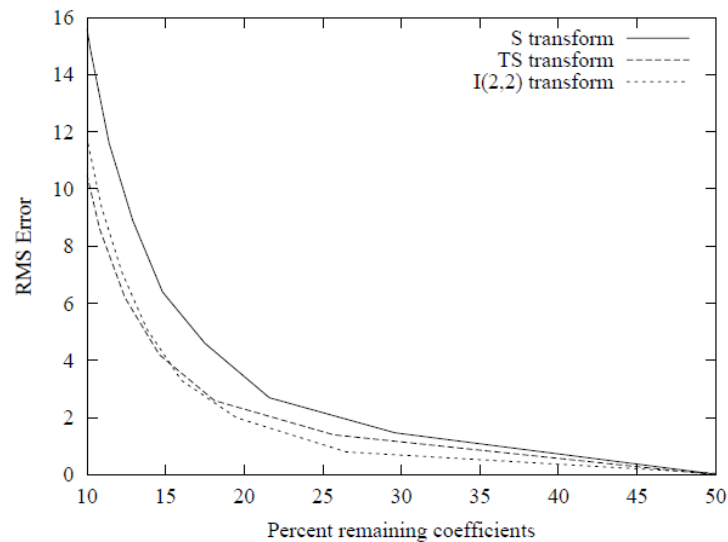


Fig. 1 : Lossy compression with the S, TS, and I(2,2) transforms: the RMS error of an image is plotted as a function of remaining non-zero coefficients.

So far we have only considered lossless wavelet domain representation of images (only coefficients that become identically zero as a result of the wavelet transform are eliminated from the representation). Lossy representations obtained by zeroing out small wavelet coefficients yield a drastic reduction in the number of remaining coefficient in return for a modest increase in RMS error, as demonstrated by the plots in Fig. 2. Such representations can be acceptable if numerical accuracy is not critical. When choosing a wavelet transform for a lossy wavelet domain representation one additional requirement must be taken into account, the graceful degradation in visual quality of the image. In this respect we found the slower biorthogonal TS transform to be superior to the S transform. More specifically, the lossy TS transform tends to produce smoother and more visually accurate results compared to the lossy S transform, which introduces blocky artifacts.

III. INTEGER WAVELET WARPING

In order to perform 3D warping in the wavelet domain, we express the warping equations as element-wise divisions between linear combinations of four matrices [8]. Let F_i denote the matrix of all the values $f_i(x, y)$, and let U and V denote the matrices containing

all of the warped u and v target coordinates. Using these matrices we rewrite equation (2) as:

$$U = \frac{F_1}{F_3} \quad V = \frac{F_2}{F_3}, \quad (1)$$

where:

$$F_i = m_{i1}A + m_{i2}B + m_{i3}C + m_{i4}D. \quad (2)$$

In the planar-to-planar warp, for example, the linear combination coefficients m_{ij} are the p_{ij} 's from equation (3), and the matrices are defined as follows:

$$A = [x]_{x,y} \quad B = [y]_{x,y} \quad C = [1]_{x,y} \quad D = [\delta(x,y)]_{x,y} \quad (3)$$

Thus, the matrix A is simply a linear ramp, increasing from left to right; all of its rows are the same vector $[0, 1, \dots, n-1]$. Similarly, the matrix B is a linear ramp, and all of its columns are the same vector. The matrix C is constant. The wavelet transform of these matrices is extremely sparse, and the efficiency of our wavelet warping algorithm stems from this sparse representation [9]. In the cylindrical-to-planar case the matrices are slightly more complicated:

$$\begin{aligned} A &= [\sin(2\pi x/w)]_{x,y} & B &= [\cos(2\pi x/w)]_{x,y} \\ C &= [y_0 + (y_1 - y_0)y/h]_{x,y} & D &= [\delta(x,y)]_{x,y} \end{aligned} \quad (4)$$

Still, note that each of the matrices A and B is a function of a single variable x , which means that in each of these two matrices all of the rows are equal. Similarly, C is a function of y , and therefore all of the columns are equal.

$$\begin{aligned} A &= [\sin(2\pi x/w) \sin(\pi y/h)]_{x,y} & B &= [\cos(2\pi x/w) \sin(\pi y/h)]_{x,y} \\ C &= [\cos(\pi y/h)]_{x,y} & D &= [\delta(x,y)]_{x,y} \end{aligned} \quad (5)$$

In this case only C is a function of a single variable y , and therefore all of the columns are equal. The wavelet warping operation consists of three steps: (i) computation of linear combinations (equation (6)), (ii)

Both the standard cylindrical-to-planar warp and our wavelet warping algorithm exploit this structure to save computations [10]. Finally, in the spherical-to-planar case the matrices are:

reconstruction, and (iii) clipping and element-wise divide. The first step is carried out in the wavelet domain. Thus, following equation (1), we compute the matrices F_i as follows:

$$\begin{aligned} F_i &= T^{-1}T(m_{i1}A + m_{i2}B + m_{i3}C + m_{i4}D) \\ &= T^{-1}(m_{i1}T(A) + m_{i2}T(B) + m_{i3}T(C) + m_{i4}T(D)) \end{aligned} \quad (6)$$

The matrices A , B , and C depend only on the type of warp (planar, cylindrical, or spherical), and are independent of the reference or the target images.

Consequently, $T(A)$, $T(B)$, $T(C)$ are precomputed once for each type of warp, and then reused for all warping operations [11]. The matrix D , which is the disparity image of the reference view, is

independent of the target view, and $T(D)$ is precomputed once for each reference view. The scalars m_{ij} are dependent upon both the reference and the target views, and are calculated once for each target view, the same as in a standard warp. The disparity values in D , as well as the entries of A , B , and C (in the

cylindrical and spherical cases), contain floating point values. These values are first mapped into an appropriate integer range, since our implementation uses an integer wavelet transform.

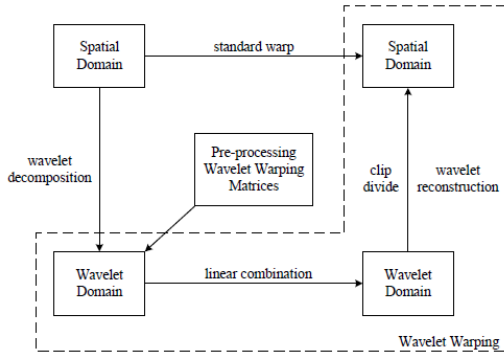


Fig. 2: Standard warp vs. wavelet warp

IV. IMPLEMENTAION OF WAVELET TRANSFORM

There are two requirements that a suitable wavelet transform should satisfy: (i) the transforms $T(A)$, $T(B)$, $T(C)$, and $T(D)$ should be sparse; (ii) the reconstructions (inverse wavelet transforms) should be fast to compute. Based on the experiments reported before, we chose a slightly modified version of the second-order interpolating wavelet transform, $I(2, 2)$. The modification consists in omitting the update phase of the lifting scheme. The resulting transform requires $83n^2$ operations to decompose an $n \times n$ matrix using the 2D nonstandard wavelet transform. The wavelet coefficients of this transform measure the extent to which the original signal fails to be linear. In the case of a planar warp, the matrices A and B are simply linear ramps and matrix C is constant (eq. 7)). Consequently,

the transforms $T(A)$ and $T(B)$ consist of two non-zero coefficients each, and $T(C)$ consists of a single non-zero coefficient. Note that this is lossless compression of the three matrices, they can be reconstructed exactly from these sparse transforms. In the case of a cylindrical warp (eq. (8)) the transforms $T(A)$ and $T(B)$ have fewer than $19n^2$ nonzero coefficients each, while $T(C)$ has two non-zero coefficients. In the case of a spherical warp (eq. (9)) the transforms $T(A)$, $T(B)$ and $T(C)$ have fewer than $19n^2$ non-zero coefficients each. Once again, the compression of the matrices is lossless. As for the disparities matrix D , the number of non-zero coefficients depends, of course, on the reference image. In our experiments, roughly one third of $T(D)$ coefficients were non-zero. Although the number of non-zero coefficients can be decreased further by lossy wavelet compression, it is not beneficial to do so. As we shall see in the next section, the computational bottleneck of wavelet warping lies in the reconstruction stage. A slight reduction in the number of coefficients does not significantly improve performance, while a more drastic truncation causes errors in the mapping, resulting in visible artifacts.

V. EMPIRICAL RESULTS

We have implemented our wavelet warping algorithm, as well as the standard warps: incremental planar-to-planar, LUT-based cylindrical-to-planar and spherical-to-planar, with the optimizations mentioned earlier. The algorithms were implemented in Java. All of the results reported in this paper were measured on a 3.0 GHz Pentium Dual Core processor. In all our comparisons we measured the entire warping time at full resolution, including reconstruction, clipping, and the divisions by the homogeneous coordinate. The averaged performance of the different warping algorithms (in frames per second) is summarized in Table 4.

Tab.3: Measured performance (frames per second) of standard vs. wavelet warp in the planar, cylindrical, and spherical cases.

Type of warp (number of pixels warped)	Standard warp	Wavelet warp
Planar (512 x 512)	6.5	7
Cylindrical (512 x 256)	12	15
Spherical (512 x 256)	7.7	14

As predicted by our analysis, we found wavelet warping to be roughly as fast as the standard algorithm in the planar case and slightly faster (up to 25 percent) in the cylindrical case. Note that in the planar case the reference image has twice as many pixels as in the cylindrical case. This is the reason that the number of warps per second in the first row of the table is smaller almost by a factor of two. As expected, in the spherical case, wavelet warping outperforms the standard algorithm by a factor of roughly 1.8.

VI. CONCLUSIONS

We have presented a simple way of computing 3D image warping in the wavelet domain. We have demonstrated both analytically and experimentally that performing these operations in the wavelet domain is in many cases faster than their direct computation. Furthermore, wavelet domain operations enable progressive and multi-resolution computations, as well as space and frequency locality. We have demonstrated

our approach both on still images and on image sequences. To extend and improve our approach, we would develop an adaptive multiresolution scheme, which would allow operating upon different regions of an image at different resolutions.

REFERENCES REFERENCES REFERENCIAS

1. Shenchang E. C. (2005) QuickTime VR — an image-based approach to virtual environment navigation. In Computer Graphics Proceedings, Annual Conference Series (Proc. SIGGRAPH '95), 29–38.
2. McMillan L., Bishop G. (2005) Plenoptic modeling: An image-based rendering system. In Computer Graphics Proceedings, Annual Conference Series (Proc. SIGGRAPH '95), 39–46.
3. Calderbank R., Daubechies I., Sweldens W., Yeo B.L. (2008) Wavelet transforms that map integers to integers. *Appl. Comput. Harmon. Anal.*, 5(3):332–369.
4. Dally W. J., McMillan L., Bishop G., Fuchs H. (2006) The delta tree: An object centered approach to image-based rendering. MIT AI Lab Technical Memo 1604, MIT. 234-245.
5. Debevec P., Hawkins T., Tchou C., Duiker H. P., Sarokin W., Sagar M. (2000) Acquiring the reflectance field of a human face. *Siggraph 2000, Computer Graphics Proceedings, AddisonWesley Longman, Los Angeles, USA*, 145–156.
6. Sweldens W. (2007) The lifting scheme: A construction of second generation wavelets. *SIAM Journal of Mathematical Analysis*, 29(2), 511–546.
7. DeVore R. A., Jawerth B., Lucier B. J. (2002). Image compression through wavelet transform coding. *IEEE Transactions on Information Theory*, 38(2 (Part II)):719–746.
8. Dutilleul P. (2008). *Wavelets: Time-Frequency Methods and Phase Space, Inverse Problems and Theoretical Imaging*, pages 298–304, Berlin. Springer-Verlag.
9. Guo H., Burrus C.S. (2006) Convolution using the undecimated discrete wavelet transform. In *Proc. Int. Conf. Acoust., Speech, Signal Processing (ICASSP-96)*, Atlanta, GA.
10. Ramamoorthi R., Hanrahan P. (2001). An efficient representation for irradiance environment maps. In *Siggraph 2001, Computer Graphics Proceedings, Los Angeles*.
11. Said A., Pearlman W.A. (2006). An image multiresolution representation for lossless and lossy image compression. *IEEE. Trans. Image Processing*, 5(9):1303–1310.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Representing Aspect Model as Graph Transformation

By Vishal Verma, Ashok Kumar

Kurukshetra University, Kurukshetra

Abstract - In this paper we discussed a new method for representing aspect models. This method uses the basics of UML to devise a new way for specifying the model level aspects and transformations among them. The resultant model is effective from both expression and scaling point of view. The work in this paper is based on assumed transaction processing system in a bank.

Keywords : Aspect - Oriented, Graph Transformation, UML.

GJCST Classification : D.2.9



Strictly as per the compliance and regulations of:



Representing Aspect Model as Graph Transformation

Vishal Verma^a, Ashok Kumar^a

Abstract - In this paper we discussed a new method for representing aspect models. This method uses the basics of UML to devise a new way for specifying the model level aspects and transformations among them. The resultant model is effective from both expression and scaling point of view. The work in this paper is based on assumed transaction processing system in a bank.

Keywords : *Aspect - Oriented, Graph Transformation, UML.*

I. INTRODUCTION

In most of the software development techniques identification and presentation of aspects is done only at some specific levels which pose constraints on the designer and developers to follow a predefined pattern/steps for development process. In this method we try to develop a technique which can be used for representing and composing aspect at any level of software development. With the advent of new techniques for software development it is quite common and natural, that aspect can occur during any of the development phase i.e. requirement [1], analysis [8] and design [12]. Aspect if modularized during software modeling can leads to a clear boundary among aspects and concerns and they become more maintainable, understandable and organize-able within the model. On the other hand if aspect modules are composed with the development of base module then it helps to fully understand and analyze the model with aspects, and any ambiguity, conflicts and omissions can be avoided. Hence, the mechanism used for specification of aspect at the modeling level must be complemented with mechanism used for composition, that weave the aspect model into base model. Lack of expression and scalability are the major problems faced by the researchers for development of mechanism like this. Composition at the modeling level can be extremely rich in nature [14]. Existing models do not provide support for expressing the richness in compositions. However, increase in the degree of expressiveness can lead to the problem of scalability because a large effort is required by the developers to specify the composition. The method discussed in this paper is capable to handle the problem of scalability and expressiveness and the result

of this paper is a practical technique that can be used for defining and composing aspect oriented model for best modeling purpose.

The method used in this model is based on two basic technology i.e. Role Based Meta Modeling language [2] and graph transformation [3, 5]. Role Based Meta Modeling language provides a precise, simple graphical means for specifying a model level aspect in a way that is consistent with UML [13]. It is used for modeling the structural part of security aspects [6] as well as model behavioral UML aspects [8]. The base problem faced while using RBML is that they do not scale up to marks since a lot of effort is required to specify the cross cuts among the core modes. Our discussed method shows clearly the reduction in level of effort to be done for models. Transformations using graphs have been applied in a number of problems related to the software engineering and to the problem of merging of different systems together [9], but in none of the implementations it has been categorically addressed how to apply them, in general way, to handle the aspect at any level of UML modeling. The aim of this paper is to combine together the RBML and graph transformations to achieve

- a) General implementation of UML based aspect modeling and composition at any stage of abstraction.
- b) To implement the proper scalability of aspect composition.

This paper illustrates the approach with an assumed transaction execution system based closely on an existing application used by banks.

II. MODEL LEVEL ASPECTS

Aspect oriented models are models which represent the cross cut, points cut and concern in a well arranged manner along with aspects. From the view point of problem discussed in this paper it can be defined as a model that crosscuts other model at the same level of abstraction. Here the words "same level of abstraction " plays very important role i.e. a model is considered to be an aspect if it crosscut the other model of same interactions e.g. if requirement cut requirement model, requirement artifact cuts requirement artifact only then they are considered as aspects. In particular case a use case may not be aspect. Although a use case is suppose to always cut

*Author ^a : Department Of Computer Science and Application, Kurukshetra University Post Graduate Regional Centre, Jind.
E-mail : vishal.verma@kuk.ac.in*

Author ^a : Department Of Computer Science and Application, Kurukshetra University, Kurukshetra.

across multiple implantations module, it is only considered to be an aspect if it cuts across other use case.

Discussion in this paper is restricted to the definition of an aspect oriented model with a condition that a model is an aspect only if it crosscuts other model built with same perspective e.g. any model which is build for global interpretation of interaction cannot cut a model build with local interpretation of interaction and hence is not at same level of abstraction but they have different perspective- local and global. These types of models are not considered in this paper.

a) Representation of Aspect in Role Based Modeling Language

Role Based Modeling Language [2] is used in this paper to represent the Aspect-Oriented Modules. This language is further complemented by France et al [10]. RBML is considered as a special case of UML Meta model in which each element of RBML is treated as a role. It is also considered that a role is a constraint of a UML Meta class with a set of optional properties that any element must possess. Because RBML is considered as a special case of UML hence each UML diagram must have a corresponding RBML diagram in which model elements are roles e.g. state roles and transitions of RBML represents a generic state that can be made concrete by assigning it to a concrete role. Proper care is to be taken that only those model elements which satisfy the property of a role should be treated as a role. RBML model defines a generic model that can be instantiated in many ways by assigning elements to its entire role [14]. Any UML model is said to conform to a RBML model if there is a valid argument of elements in the UML model to the roles of RBML model [14].

RBML model is used to formalize the design pattern [2] and to represent model level aspect [4]. This was extended to behavioral aspect in [8] and [7]. As per the original definition in [10] all RBML model elements must be roles i.e. they are Meta – level elements. As per [8] for representing aspects it is useful to allow object-level elements in RBML as well. The result is an extended RBML, represented by eRBML, in which an element may be Meta level or object level element [14]. Fig.1 shows the sequence of aspect in eRBML. It clearly shows that whenever the user get ack of failed transaction the HOST itself record the status in STATUS file and at the same time shut down the USER side as well. Fig 1 shows the combination of object level elements meta-level role together in one go. This type of combination is preferred since status like objects are remains unchanged and their relative updation dependent on the varying values of roles only.

- i. *Instantiation: it means to assign some concrete*

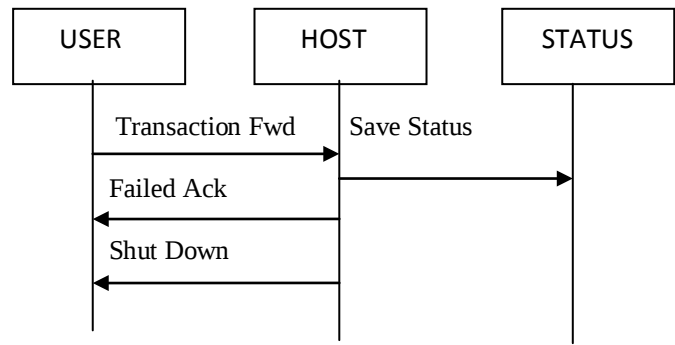


Fig.1 : Handling of Failed Transactions

values to the elements or one –to –one mapping from role to model elements. In context of eRBML each aspect model must be instantiated before it can be composed with a base model. Instantiation is basically used to define what the aspect should like in context of a particular application i.e. the aspect is identified and specialized to a context. Fig.2 shows another example of how aspects cross cuts each other. Sequence diagram in Fig. 2 is taken as base for further discussion and is part of our case study in coming sections. From fig. 2 it is enough to conclude that there is a controller which keeps control of accessing request from user and sending it to the server for processing. Controlling all aspect of transaction is the sole responsibility of the

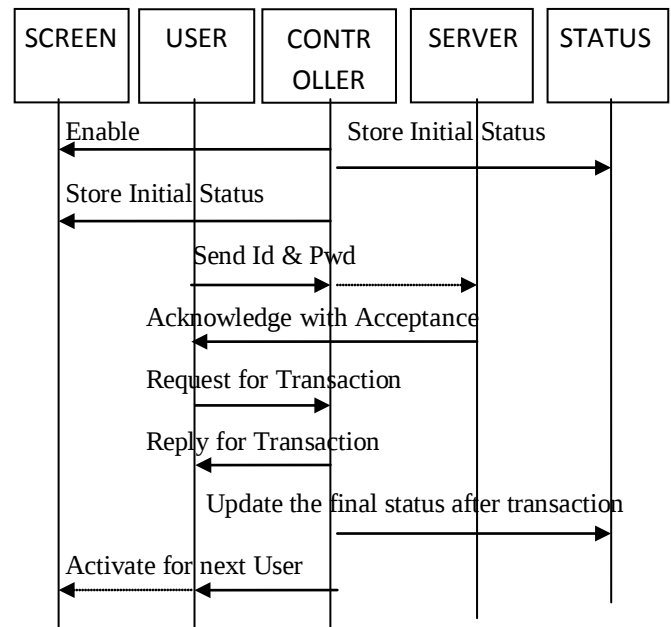


Fig. 2 : Base Sequence Diagram for user Transactions

controller, it also provide the necessary GUI for processing. Failure handling is not considered as part of this discussion. Instantiation is used to propose the aspect for composition with the base. In the example discussed here following instantiations are specified by the modeler: | USER -> CONTROLLER| CONTROLLER-> SERVER and failure are not considered.

Fig.3 shows the result of composition of Fig. 1 and Fig. 2 and instantiate the aspect with base modeler. Message to deal with intermediate REQ is included and is provided as an alternate execution path using UML 2.0 alt interaction fragment.

ii. Conformance

A UML model is said to conform w.r.t. eRBML if their exist an instantiation of eRBML model in a way that all elements of instantiated eRBML are present in UML model along with existence of constraints. The constraints are suppose to include the message ordering, sequence diagram, transition ordering and additionally specified properties of eRBML roles with respect to an UML model there may exist any number of different eRBML models that conforms dependency upon the availability of additional transaction. There may be any number of additional transactions that exist in between starting and closing transaction i.e. intermediate transactions. Hence conformance should be considered as a type of refinement.

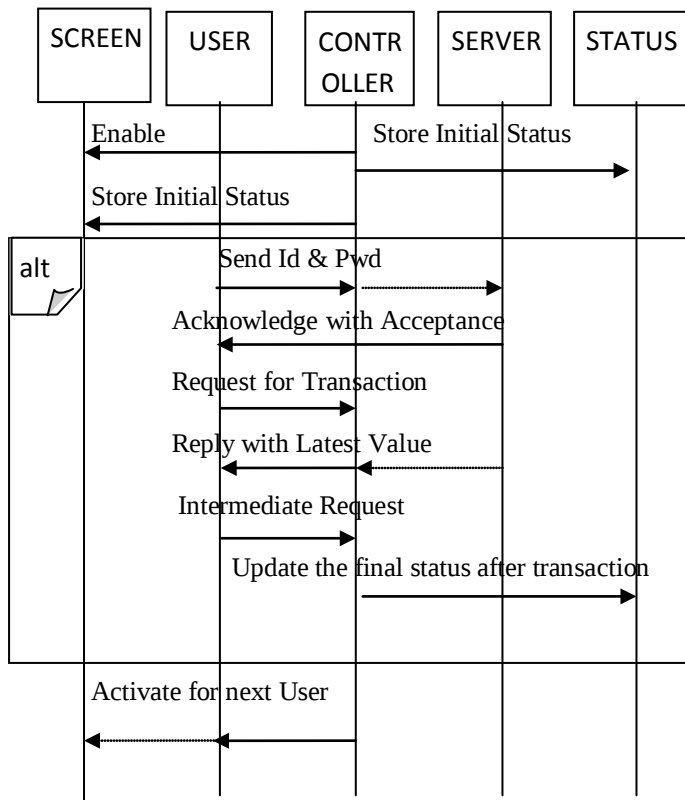


Fig. 3 : Composite of Fig.1 and Fig. 2

b) Aspect Composition

Model level aspect can be specified in a well defined way in eRBML with respect to the aspects of any UML diagram. As it is necessary to represent the model aspect in a modular fashion, composing the model aspect with base model is also important. Here we are comparing the composition approaches of France et al [11] and Whittle [8] to identify the limitations

of existing method to model. First one out of these two approaches use templates to represent aspects. Instantiation of eRBML aspect before composition is mandatory in both of the approaches, though both of the approaches [11] and [8] have different way of implantation of composition.

We have discussed a single method for composition in Fig. 3. This figure though uses simple technique, yet there are many alternates by using which composition could be done. In fig.3 the intermediate transaction is introduced which can be placed at different level of execution and can produce different compositions accordingly. Although it is simple in nature, it may be not be suitable in many cases. Common limitations of this method is that it is not able to specify the fact that how the aspect messages should be interleaved with the base model, or to specify that the aspect messages define a sequence executed in parallel with the base model message. As an alternative there are many possible ways that composition could occur. Challenging part to find a way for specifying composition that admits a high degree of expressiveness, with minimum effort to be applied for modeling. To find the response on expressiveness and scalability below we compare the techniques discussed in [11] and [8].

The method discussed in [11] allows the modeler to describe the composition directive that finally tailors the tailor algorithm. These composition directives permit the user to specify the aspect message interleaved with base or as an alternative or to run in parallel of it. "Addition", "deletion", and "move", statements are supposed to be used as directives to make the composed model. To merge the base and instantiated model first of all their elements with the same name are merged together. On completion of merging of elements with same name in first go directives are tailored to drive the exact form of composition. This all method of tailoring demands a lot pressure on modeler. Manual composition in this way demands a composition to be implemented by first applying the directive in each and every model's element also in the base model. This type of procedure can't be scaled at all.

In contrast to this the method discussed in [8] is at higher level of abstraction. In this method the composition operators are used instead of composition directives. Specifically AND, OR and IN operators are used. AND operator is used to interleave the base model with aspect model. OR can be used to provide the alternative sequence among base and aspect model. IN has some special use and it is used to insert the aspect message in any base sequence. Operators used in this approach offer a high level view of composing aspect models. This approach is more suitable and easy then the one discussed in [11] since user is not required to work with element by element manner. The only drawback of this approach is limited

number of operators for performing all desired tasks.

III. NEW METHOD OF COMPOSITION

Techniques discussed in sec 2 have limitations with respect to scalability. In that technique it is the sole responsibility of modeler to specify a set of role instantiations for each aspect and for each base model that is cross cut by aspect. It is obvious that for large aspect more number of instantiations is required to be supplied. From fig. 3 it is clear that all the instantiations are given in non graphical and in low level format that are time consuming to understand and ultimately make the maintenance of the model more difficult for the user. In comparison the method discussed in this paper provide a clean and clear way of separation of aspects and the base model. The newly discussed technique provides a new way of representing and composing model aspects in a way that maintains aspect modularity along with scalability. Basically graph transformation rules are used for representing composition and is represented by a rule ($L \rightarrow R$) bearing left hand side and right hand side. Left side is responsible for keeping points where the aspect should be applied and right side keeps the aspects in it.

a) Aspect as Graph Transformation

A graph transformation discussed in [5] is a rule represented by r and has L as left hand side and R as right hand side. Rule r is supposed to be applied on a graph G and the process of applying r finds a graph homomorphism, h , [5] from L to G and replacing $h(L)$ in G with $h(R)$. To avoid any kind of unreferenced edges i.e. edges with missing resources or target node - $L(R)$ is applied into G in such a way that all edges connected to a removed node in $h(L)$ are reconnected to a replacement node in $h(R)$.

UML diagram can be represented in the form of a graph because it is defined by the UML meta-model which is a graph where the nodes are Meta classes and the edges are meta-relationships [13]. Hence it is possible to represent transformation over UML model as graph transformation. Particularly we see composition of an aspect model with a base UML model as a graph transformation LHS and RHS both are eRBML models. As above L side specifies the points where the aspect should be applied and the R side specifies the crosscutting structure/behavior that should be inserted at those points.

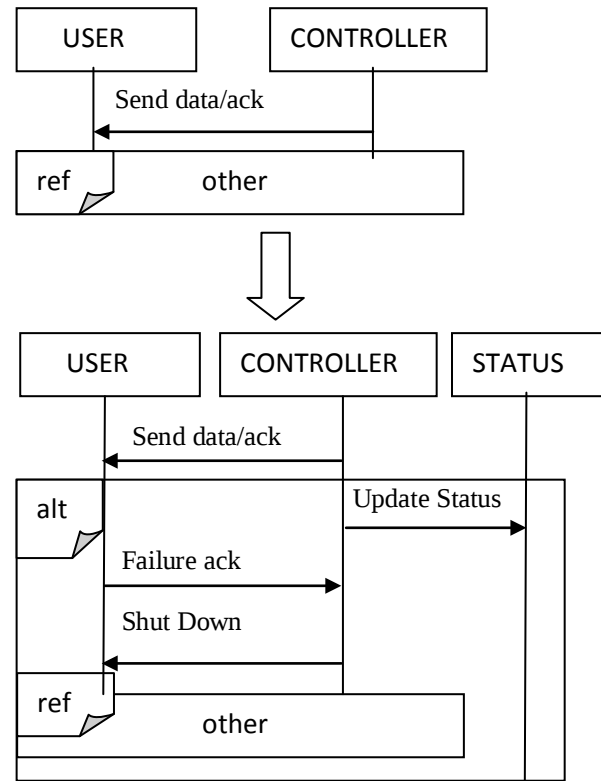


Fig 4 : Composition as Graph Transformation and handling Failure Aspects

Fig 4 represents how aspects from fig. 1 can be represented by using graph transformation. It shows two parts of aspects definition. Left side defines the aspects itself and Right side defines the composition strategy. On applying composition it would become possible to deal with message for future that is inserted as an alternative sequence after all instances of send data/ack, a message sent from CONTROLLER to USER. The approach used here helps to define the expressiveness and scalability related to composition in an easy way. It becomes clear from fig. 4 that it becomes possible to keep a complete separation of the aspects and its composition strategy. It helps to reuse of the aspects and application of the same aspects with different purpose and different composition strategy. This technique is a fully expressive way of defining composition strategy – as one is shown in Fig. 4(it is one other alternative may be used). This strategy uses the number of instantiations required to design a model. In the example discussed in Fig . 4 only one instantiation is required to be provided by the modeler i.e. failure ack. Rest all roles can be instantiated by graph matching against a base mode, the left side of the graph transformation is required to be matched with base model instantiating USER, SERVER , send data/ack (only failure ack is required to be instantiated). In fig.4 the UML 2.0 ref fragment is used to specify the placeholder for a sequence of messages in the base. This is an easy way to match against a message sequence whose position in the composed model can then be specified exactly on the R.H.S of the

transformation. At the time of defining the graph transformation w.r.t eRBML, the base definition must be modified. When matching against the Left side of the transformation rule r , it is mandatory to discuss the instantiation for the role elements. In this context the base definition is modified as the graph transformation applies to a UML model if and only if the Left side of transformation has a graph match i.e. module conformance exist there. The method in terms of scalability and expressiveness can be defined as:

i. *Scalability*

Main limitation of the scalability of all aspect approaches based on RBML is that the modeler must instantiate the role elements for each base model crosscut by the aspects. Use of graph transformations reduce this effort because instantiating the role elements become automated to some extent. Instantiation place a need to find a base model over which graph transformation can be applied- i.e. finding a match for left hand side of the transformation rule. As per above discussion we apply the module conformance while applying an aspect i.e. while working with the eRBML model, R (rule), match a UML model say U , modulo conformance if and only if there is an instantiation of the role element $\emptyset$, such that $\emptyset(R)$ conforms to model U . as is clear from Fig. 4 it has modulo conformance with the base model and hence problem of scalability is managed well by graph transformations.

ii. *Expressiveness*

As shown in Fig. 4 and sec 3.1 the Right side of the graph transformation rule, r , defines the manner in which the aspect cross cuts a base model. Since Right side is a model in itself, it completely reflects the expressiveness and how the cross cutting is defined. Here the aspect messages can be defined as an alternative to a base message or messages, as interleaved with the base message, accessing in parallel with the base message or any other combination of above discussed alternatives. The composition operator discussed in sec 2.2.2 can be defined as special case but graph transformation allows any combination of these operators to be specified, or needed for new operators to be specified. The composition directive in sec 2.2.1 are subsumed by the graph transformation approach because there is no longer any need to tailor the aspect composition algorithm to add, delete or remove elements - these modifications are rather defined explicitly in the Right side of the transformation.

IV. CASE STUDY

For the purpose of doing the case study of the expected system to be developed, we assume a system in general which is responsible for processing of transactions raised by user in terms of bank transactions. Every user of bank is supposed to execute a set of queries (may be predefined) for completion of desired tasks. We assume a system for study in which

each user is required to first authenticate him/herself for executing other transactions. After authentication use is provided with a GUI by using which rest all requests can be processed. Some of the simplest form of query is deposit and withdraw of amount and to get a balance or mini statement from the bank. In all of this type of queries a proper integrity among user interface window and ATM machine is mandatory i.e. any transaction which affect the balance in the account must be effective at all place and do the final status change at some common location. These type of queries are expected to be executed from ATM machine, from online based banking system, mobile based banking system or from a window in a bank's office where a bank officer is supposed to execute desired on verification of credentials from user. Important among all these alternatives of query execution is that they must do final status change at common location which is accessed by all means of query execution and all the time latest updated value must be available at that location i.e. SERVER. Data integrity is clearly an important issue to be maintained in design of this type of systems. Any user who is permitted to use his account by a number of means is dependent on one central location i.e. SERVER for latest updated values. The design of this system is done by using the two - phase commit protocol for maintaining serializability among the transactions to keep the commonly used values updated at all the time. The stress here is on the application implementation of protocol not on its practical details and is embedded in the working of CONTROLLER. Following we are showing the embedding of protocol in the CONTROLLER (core) functionality and how the protocol is implemented via aspects. This implementation is easily readable and hence any desired changes in integrity are easily implementable. The design is done in UML by keeping the dynamic nature of the design.

Importantly two situations demands the close look upon the updated data i.e. first when a user is accessing the account by bank window and at the same time accessing via the mobile banking services and second when transaction through ATM is under process and at the same time mobile banking transaction is executing. Both of the situations demand the very proper execution of two phase commit protocol.

Fig. 5 and Fig. 6 shows the base sequence execution of transaction models corresponding to the execution of query's from WINDOW and MOBILE at a time and from ATM and MOBILE at a time. In second discussed scene the execution of query and updation in final value may be delayed for some time because updation done through ATM may take some time for final updation in the system. In fig. 6 we are introducing a new syntax (all) used for processing the number of transactions together finally at the common server location. Transfer of amount among the inter bank accounts processed in this way cause delayed

execution. The intermediate results may get stored in **CONTROLLER** and are finally updated to the **SERVER**. Here the two phase commit protocol is not modeled as part of the core

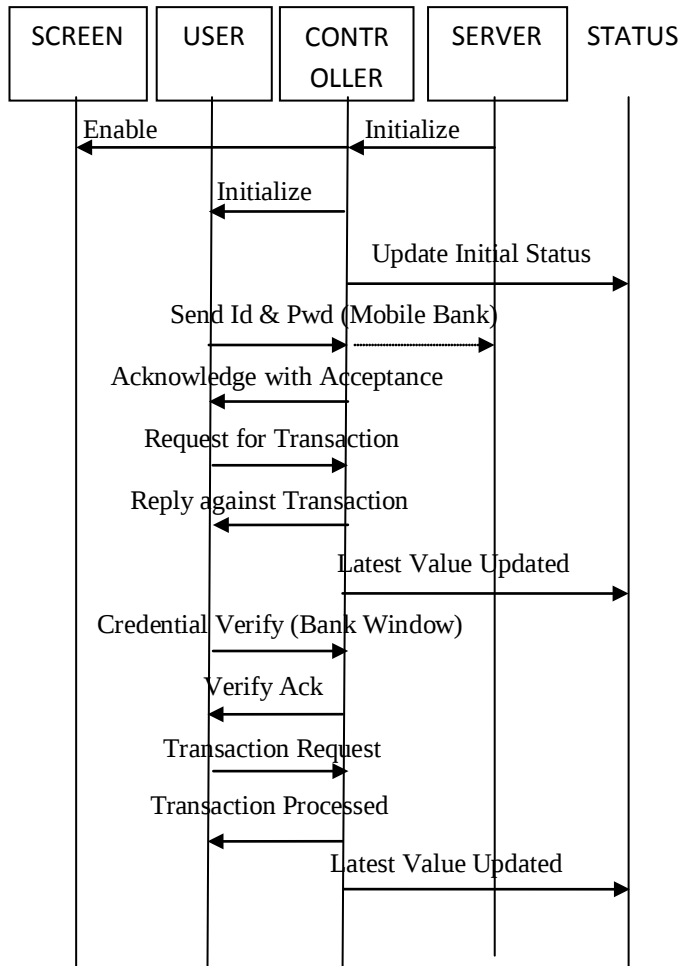


Fig. 5 : Maintain Serializaibility among parallel execution of Window and Mobile banking (single user)

functionality. Rather it is modeled separately for easy modification if needed. In Fig.5 and Fig. 6 every time the trigger of transaction is initialized by initializing both the server as well as client by **CONTROLLER**. Then first of all data (initial) is updated at sever and first **GUI** is provided to the client. In steps proceeding further the **ID** and **PWD** is submitted from **USER** to the **SEVER** and on receiving the **ACK (POSITIVE)** further transactions are processed.

Two phase commit protocol is modeled and shown in Fig. 7. It is build by considering the aspectual view of transaction and keeping them in sequence in an **eRBML**. Aspects are used to define a general pattern of communication to be used by

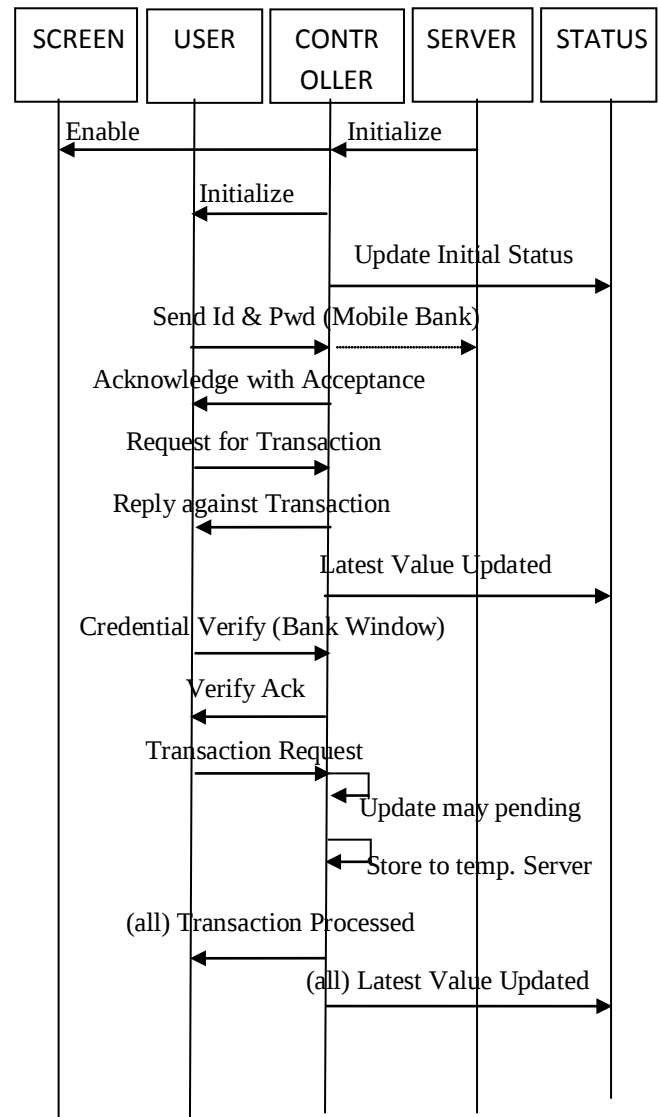


Fig. 6 : Maintain Serializaibility among parallel execution of Window and Mobile banking (single user) (Updated Fig. 5)

protocols and are easy to modify for reusability at any stage of application development. Two identified elements in Fig. 7 are **USER** and **COMMIT SERVER**. Interaction among the two is given in the form of message role so that it can be instantiated whenever required with any specific message names. Important implication of Fig. 7 is that it commit any of the transaction only if both the **USER** as well as **COMMIT SEVER** agrees on the transaction. This all is modified and is shown in Fig. 8

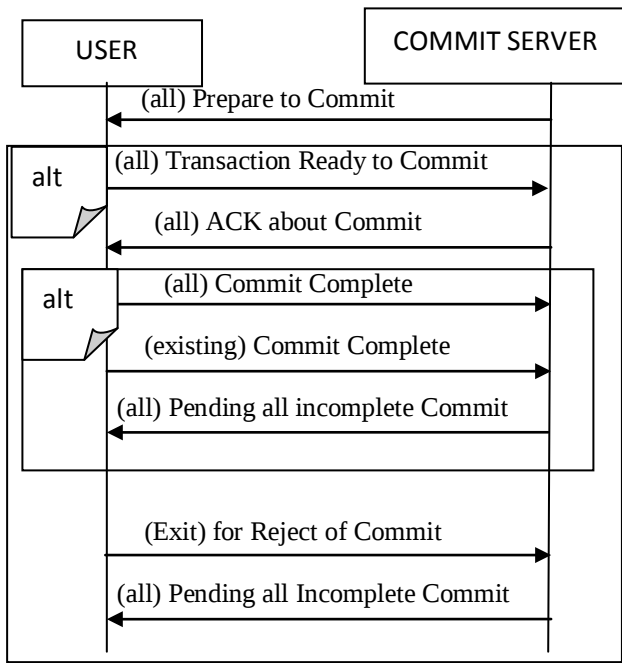


Fig. 7 : Two Phase Commit Protocol

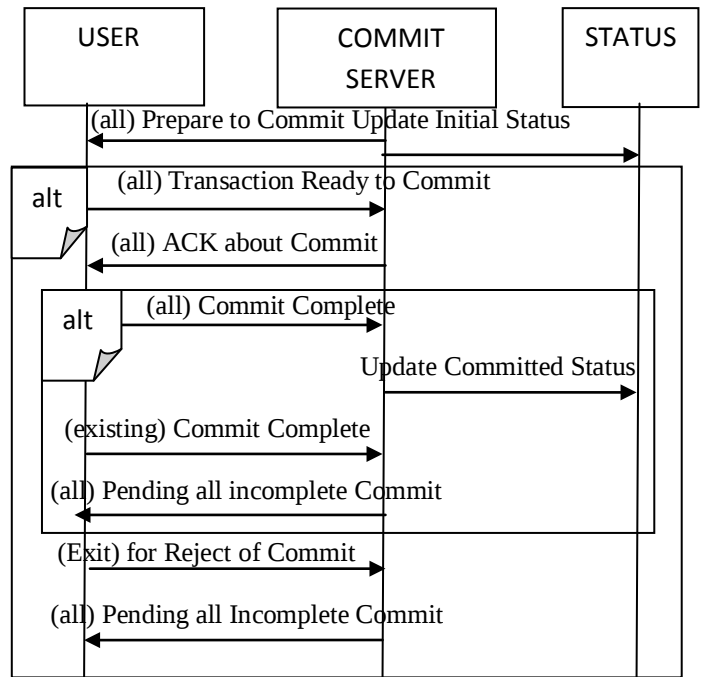


Fig. 8 : Updated Two Phase Commit Protocol

Fig. 9 and Fig. 10 shows the left side and Right side of graph transformation for refined aspect discussed in Fig. 6. Important to note here is that the Left side says that we have to apply the aspect at the points at which prepare for commit will appear and the same should preceded by Initialize message. The enable is true for both of first and second scene. Here it is possible to process the step by step manner or to execute a separate algorithm for execution of transactions. Right side of graph is shown in Fig. 10. In this fig other messages are included to take into consideration all or any kind of transaction which not be used in general by all user but is expected to execute in some special case only. The messages are supposed to be executed only if the reply from the two phase commit protocol is true. Two phase commit protocol is able to reply true or false depending on the execution of transactions. The base and aspect model are composed in such a way that match for all other messages is done only after point of successful commitment. The use of existing method discussed in [11] and [8] are not able to specify the conditions. The method presented by [8] may allow the weaving in the way which we want to describe. In method discussed in [11] it is needed to specify a list of composition directives that give the instructions to composition algorithm where to place the messages matching with other messages. Hence the messages

discussed in [11] and [8] are not appropriate for presenting the directives in easy way and are time consuming and error-prone too. In comparison to these two methods the graph transformation is an easy graphical method to specify the directives. In the method suggested in this paper it is very easy to place any additional messages anywhere in the Right side of the graph transformation rule. Along with is also possible to specify a different composition way to simplify/modifying the Right side of the Rule.

V. CONCLUSION AND RELATED WORK

All of the existing approaches used to identify, compose and represent aspect at various level of software development faces a number of limitations especially the problem of scalability. The approach discussed in this paper for representing aspect at any level of software development using the UML methodology based on role modeling language. Various level of hierarchy are used to structure aspects and their possible instances. The problem of scalability is sorted out in this method since graph transformation allow the matching at any level of development and it automatically compose aspects along with the problem of expressiveness is also sorted out as use

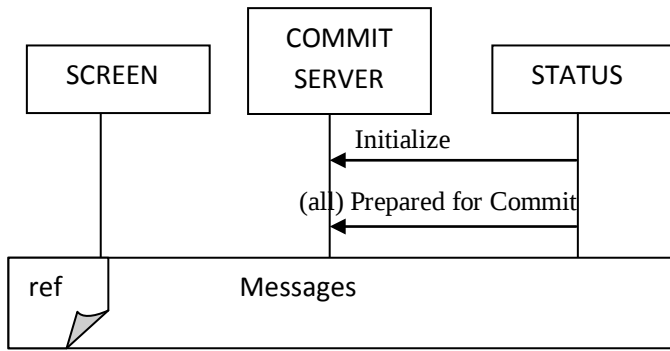


Fig. 9: Left Hand Side of Graph Transformation

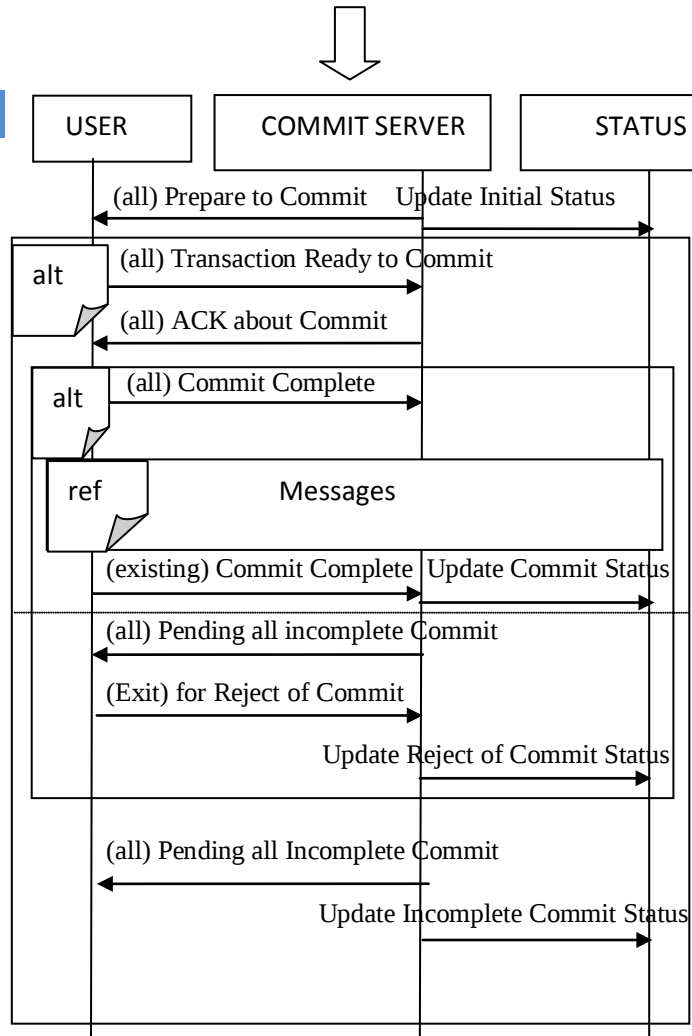


Fig. 10: Right Hand Side of Graph Transformation

of graphical method in terms of graph transformation expresses all implementations. The validity of approach is reflected through its use on bank's transactions.

The approach discussed in this method is more close to syntactic implementations a lot of modifications can be done in terms of syntax related issues so that immediate implementations in programming language can be done. The modification to resolve the conflict among the aspects can also be done. The matching

process discussed in this paper is also open to be modified. The modification in terms of forward and backward movements on matching at any level of transformation can be done.

REFERENCES REFERENCES REFERENCIAS

1. A. Rashid, A. Moreira and J. Araujo, (2003) "Modularisation and Composition of Aspectual Requirements". AOSD 2003 Boston, USA, ACM Press, pp 11-20, March 2003.
2. Dae-Kyoo Kim, Robert France, Sudipto Ghosh and Eunjee Song, "A Role- Based Metamodeling Approach to Sepecifying Design Patterns", COMPSAC 2003, Dallas, Texas, 2003.
3. D-K. Kim, "A Metamodeling Approach to Sepcifying Patterns", PhD Thesis, Colorado State Univreisty, 2004.
4. G. Georg and R. France, "UML Aspect Specification using Role Modles", OOIS, France, Lecture Notes in Computer Science, Springer, Vol. 2425, pp 186-191, September 2002.
5. G. Rozenberg, editor. "Handbook of Graph Grammers and Computing by Graph Transformation, Volume 1: Foundations." World Scientific, 1997.
6. I.Ray, N. Li, R. B. France and D-K.Kim, "Using UML to Visualize Role-based Access Control Constraints." SACMAT 2004: 115-124.
7. J. Araujo, J. Whittle and D. K,Kim, "Modeling and Computing Scenario- Based Requirements with Aspects." RE 2004,Kyoto, Japan, IEEE CS Press, 2004.
8. J. Whittle and J. Araujo, "Scenario Modeling with Aspects", IEEE Proceedings Software, Vol 151(4) pp. 157-172,2004.
9. M. Goedicke,B. Enders, T. Meyer, G. Taentzer, "ViewPoint-Oriented Software Development : Tool Support for Integrating Mutiple Prespective by Distributed Graph Transformation." TACAS 2000; 43-47.
10. R. France, /D-K,Kim,S.Ghosh andE. Song. "A UML- Based Pattern Sepcification Technique." IEEE Transactions on Software Engineering, Vol 30(3), pp 193-206,2004.
11. R. France I. Ray,G. Georg and S. Ghosh, "An Aspect-Oriented Approach to Design Modeling", IEEE Proceedings Software, Vol 151(4), pp 174-186, August 2004.
12. S. Clarke and R. J. Walker, "Composition Patterns : An Approach to Designing Reusable Aspects". ICSE 2001, Toronto, Canada, IEEE CS Press, pp 5-14,2001.
13. UML version 2.0 Available from the Object Management Group, 2005, <http://www.omg.org>.
14. Whittle. Jon, Araujo. Joao, Moreira. Ana, "Composing Aspect Models with Graph Transformations", EA 06, Shanghai, China. 2006.



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Extended Apriori for association rule mining: Diminution based utility weightage measuring approach

By P. Laxmi, A. Poongodai, D. Sujatha

Department of CSE

Abstract - The field of Association rule mining is a dynamic area for innovation of knowledge through which uncountable procedures have been expounded. Recently, by including significant components viz. value (utility), volume of items (weight) etc, the researchers have enhanced the quality of association rule mining for industry by bringing out the association designs. In this note, a proficient methodology has been put forward based on weight factor and utility for effective digging out of important association rules. At the very beginning, a traditional Apriori algorithm has been utilized that make use of the anti-monotone property which states that if n items are recurring continuously then $n-1$ items should also recur by which the scores of weightage(W-Gain), utility(U-Gain) and diminution(D-sum), are derived at. Eventually, we derive a subset of important association rules through which EUW-Score is generated. The tentative outcome demonstrates the effectiveness of the methodology in generating high utility association rules that is profitably used for the business improvement.

Keywords : Association Rule Mining (ARM), Recurrent item set, Utility, Weightage, Apriori, Utility Gain (U-Gain), Weighted gain (W-Gain), Diminution sum (D-sum), Exact Utility Weighted Score (EUW-score).

GJCST Classification : H.2.8



EXTENDED APRIORI FOR ASSOCIATION RULE MINING DIMINUTION BASED UTILITY WEIGHTAGE MEASURING APPROACH

Strictly as per the compliance and regulations of:



Extended Apriori for association rule mining: Diminution based utility weightage measuring approach

P. Laxmi^α, A. Poongodai^Ω, D. Sujatha^β

Abstract - The field of Association rule mining is a dynamic area for innovation of knowledge through which uncountable procedures have been expounded. Recently, by including significant components viz. value (utility), volume of items (weight) etc, the researchers have enhanced the quality of association rule mining for industry by bringing out the association designs. In this note, a proficient methodology has been put forward based on weight factor and utility for effective digging out of important association rules. At the very beginning, a traditional Apriori algorithm has been utilized that make use of the anti-monotone property which states that if n items are recurring continuously then n-1 items should also recur by which the scores of weightage(W-Gain), utility(U-Gain) and diminution(D-sum), are derived at. Eventually, we derive a subset of important association rules through which EUW-Score is generated. The tentative outcome demonstrates the effectiveness of the methodology in generating high utility association rules that is profitably used for the business improvement.

Keywords : Association Rule Mining (ARM), Recurrent item set, Utility, Weightage, Apriori, Utility Gain (U-Gain), Weighted gain (W-Gain), Diminution sum (D-sum), Exact Utility Weighted Score (EUW-score).

1. INTRODUCTION

Since accessibility of huge amounts of data and knowledge is on the rise, data mining has occupied a significant place in the field of information industry [1]. Data mining is a vital part of the process of Knowledge Discovery in Databases (KDD) [22] which is a non-trivial excavation of hidden, implied, never revealed data and comparatively has a large usage [6].

In broad, data mining methods are categorized into two ways:

1. Descriptive mining

It is an account of putting forward a group of data and its characteristics in a succinct and summarized way. One of the most significant of the descriptive kind mining is Association Rule Mining (ARM) which was introduced by Agarwal et.al. [2].

Author^α : M.Tech, Department of CSE, ATRI, Parvathapur, Uppal, Hyderabad, India. E-mail : laxmi.p16@gmail.com

Author^Ω : Assoc.Prof, Department of CSE, ATRI, Parvathapur, Uppal, Hyderabad, India. E-mail : a_poongodai@yahoo.co.in

Author^β : Assoc.Prof and HOD Department of CSE, ATRI, Parvathapur, Uppal, Hyderabad, India. E-mail : sujatha.dandu@gmail.com

2. Predictive mining

This method makes to surmise the outline of the information to give some assumptions [13].

Progressively many network side scholars and researchers of computer science, particularly those who are dedicatedly working in the area of Knowledge Discovery in Data (KDD), mainly concentrate and accentuate on Association Rule Mining (ARM) [14]. Association rule mining [2, 3] has been extensively utilized to detect and unravel data mining complicated issues which involve financial troubles, and business dealings [4].

The difficulties of using the mining association rules are divided into two steps.

1. First step is to evaluate the item sets that often occur in the databases.
2. The second one is to produce the association rules. Just the once if the some item sets are found occurring often, then the production of the association rules is uncomplicated and can be accomplished in a specified time [5].

Conventional ARM algorithms judge all the data items in the similar passion by assuming that their weight-age as always 1 if they are identified or 0 if they are not identified which intangibly drives to miss some of the very functional outlines of the data [7]. In order to trounce these drawbacks of the traditional way of mining, utility mining method [9] [10] and weighted association rule mining [16] has come into existence.

Utility data mining is a latest field of development entranced in every type of utility factors in data mining procedures and accentuated to assimilate the services also called as utilities in the data mining methods [15]. Utility of any particular data or item is a reliant on the individuals and is measured in terms of aesthetic values or other expressions of individual inclination [7]. When an action in the database and its related minimum utility threshold value and a utility chart are monitored then the objective of the utility mining is to identify and determine each and every high utility item set [12]. In general pros and cons of the item set in the business is not possible to be derived by the utilization of the values that shore up the rules. So this rationally proves that utility mining can be more advantageous than the conventional association rule mining [12].

But while considering the weighted association rule mining (WARM) [16], there has been a modification of not counting the item sets that occurred in an action of the database which made a compulsion to acclimatize the conventional help to the weighted one [23]. This method also segments the consumers on the basis of their reliable nature and impending count of procurements [8]. For an illustration, one consumer may buy 15 items and some other may have 5 items at a time but the conventional association rule method considers these couple of actions in a similar way. Hence the procedure of considering the actions in a same conduct in the conventional association rule method mislays some of the vital information [8]. As a result, Weighted ARM deals mainly with the magnitude of individual substances in a database [17, 18, 19]. For a case in point, goods that has more advantages or which are under the process of endorsement are given prior able constraints in comparison to rest [20].

Incorporating the mutual characteristics (weightage and utility) for excavation of rules is treated as an addition to the weighted association rule mining which means that the data weights are most important in a particular set of actions besides it also deals with the number of possible appearances of the data in those actions. This has made a concern in categorizing the data appearances and their weights and also in detecting the prior able data which put in more to the benefits of the business [21]. By considering this Sandhu et al[28] proposed a model that identifying association rules based on UW-Score, which calculated based on characteristics weightage and utility. In their proposal they were not considering diminution occurred in contextual factors

Here, we put forward an efficient method that makes the utilization of the traditional Apriori algorithm to engender a group of association rules from a database out of which a pooled Exact Utility Weighted Score (EUW-Score) is calculated. In due course, sub values of the priority given weighted, utility and diminution considered constraints are derived on the source of the EUW-score and the tentative outcome exhibit the effectiveness.

The later part of the paper is structured as given: The section 2 explains the methods involved. The proposed methodology based on usage of mining utility-oriented association rules is explained in section 3. Section 4 concludes the paper.

II. METHODS INVOLVED

Apriori, a notable algorithm for ARM, is one of the frequently used processes of discovering an assortment of data properties which functionally is associated to bring about the data and are chiefly based on range of occurrences. But the extraction through the number of occurrences does not bring in the attention of the scholars, and to overcome this some

more measures are included in the Apriori algorithm for an efficient mining of association rules. They are:

a) *Weightage*

Unlike the general transaction database which projects the total amount of characteristics by some number, the traditional algorithms like Apriori mine association rules utilizes a binary mapped database that depicts the occurrence of the data or an item in one course of an action, thus allowing to gather and verify some good number of information related to the characteristics of the data, that results in recurrent but few number of weight-age rules. Even in an ordinary user transaction, some times the data that has a good weight-age occurs rarely, but it must also be involved in the recurrent item set. This procedure is followed in our approach, for mining a subset that has high significance.

b) *Utility*

The individual utility (Gain) of the characteristics is the subsequent measure involved in the approach to give a good standard to the ARM.

c) *Diminution*

The individual Diminution that occurs when item failed to raise the utility (Gain), which balance the utility and provide actual gain of the characteristics is the subsequent measure involved to the approach to give a good standard to the ARM. Some of the service standards in the business would be neglected in a process of mining. As these rules, when mined without these service standards will lead to a plausible loss. Those standards are attained by this method through utility measure (U-gain). Weight-age and utility measures are individually incorporated in copious researches [21, 24, 25, 26] so as to make their methods more efficient but those procedures need high capability. These procedures are effectively utilized in this methodology to extract the association rules from a database.

III. PROPOSED METHODOLOGY

Assuming D as a database having n number of transactions T and m number of attributes $I = \{i_1, i_2, \dots, i_m\}$ with positive real number weights W_i . U_i specifies the profit associated with the i attribute of utility table U with m count of utility values.

The methodology based on weight-age and utility involves some key steps like:

Step1: Extraction of the association rules from D by utilizing Apriori.

Step 2: Generation of W-gain value.

Step 3: Generation of U-gain value.

Step 4: Generation of D-sum value.

Step 5: Generation of DUW-score through W-gain and U-gain.

Step 6: Deriving the vital association rules by taking UW-score into consideration.

a) *Extraction of the Association Rules through Apriori*

To begin with, association rules are excavated from a transaction database D with n transactions. We represent the database D as:

$$D = \begin{bmatrix} T_1 \\ T_2 \\ \vdots \\ T_n \end{bmatrix} \quad (1)$$

Each transaction T in D encompasses with 'm' number of attributes $I = [i_1, i_2, \dots, i_m]$ related to it and each attribute i is symbolized by weights W_i .

To extract the rules, a typical Apriori algorithm is used in our methodology. A binary mapped database BT is applied for extracting the association rules in conventional Apriori process by which the initial database D is converted to binary mapped database BT such that it comprises of binary 0 and 1 denoting the non-existence and existence of attributes. By the succeeding equation the weights W_i are mapped onto the binary values.

$$B_r = \begin{cases} 0 & \text{if } W_i = 0 \quad \forall T_k \\ 1 & \text{if } W_i \neq 0 \quad \forall T_k \end{cases} \quad (2)$$

Consequently, an input to the Apriori algorithm [2] is produced by the binary mapped database BT for extraction of association rules which are processed in two steps of Apriori as follows:

- **Recurrent Item set Generation:** Produces min-support which signifies that each and every feasible set of attributes that comprise support value higher than a predefined threshold.
- **Association Rule Generation:** Produces min-confidence which signifies that association rules from the item sets that comprise confidence higher than a predefined threshold.

The composition of a typical association constraint is: $A \rightarrow B$, where A symbolizes the antecedent and B symbolizes the consequent and these are subset of the items in the binary mapped database, such that $A \subset I$, $B \subset I$ and $A \cap B = \emptyset$ and is construed as B co-existing if A exists. The support S and confidence C clasps the constraint $A \rightarrow B$ in the transaction database D, if item sets A and B are contained in S% and C % of the transactions. Therefore:

$$\text{Support } (A \rightarrow B) = P(A \cup B) \quad (3)$$

$$\text{Confidence } (A \rightarrow B) = P(B|A) = \text{support } (A \cup B) / \text{support } (A) \quad (4)$$

The pseudo code for the Apriori algorithm is:

```

 $I_1 = \{l \arg e \ 1 - \text{itemsets}\};$ 
for( $k = 2; I_{k-1} \neq 0; k++$ ) do begin
 $C_k = \text{apriori} - \text{gen}(I_{k-1});$  // New candidates
for all transactions  $T \in D$  do begin  $C_r = \text{subset}(C_k, T);$ 
for all candidates  $c \in C_T$  do
 $c.\text{count}++;$ 
end
end
 $I_k = \{c \in C_k \mid c.\text{count} \geq \text{min sup}\}$ 
end
Answer =  $\bigcup_k I_k;$ 
    
```

A k amount of association rules $R = [R_1, R_2, \dots, R_k]$ are produced with apriori algorithm and is sent as input to the successive part in the methodology for weight-age and utility calculation. Each attribute of k association rules of R the is determined with some measures, that is for an association rule R i of the form, $[A, B] \Rightarrow C$, where, A, B and C are considered as the attributes, the derivations U-gain, W-gain and UW-score are evaluated for each attribute A, B and C independently.

At first, a decremented arrangement is done for the produced k association rules considering their respected confidence level. The listing of rearranged association rules is specified by

$S = \{ R_1', R_2', \dots, R_k' \}, S \in R$, where $\text{conf}(R_1') \geq \text{conf}(R_2') \geq \dots \geq \text{conf}(R_k')$.

b) *Generating the value of W- gain*

At the start the initial rule R_1' is chosen from the rearranged list S and the independent attributes of R_1' are derived followed by the computation of W-gain.

Definition 1: Item weight (W_i): Item weight value W_i , is a nonnegative integer which is termed as the total magnitude evaluation of the attribute present in the transaction database D.

Definition 2: Weighted Gain (W-gain): W-gain is termed as the summation of weights of each item W_i of an attribute that is involved in each and every transaction of the database D as referred in the given equation:

$$W - \text{gain} = \sum_{i=1}^{|T|} W_i \quad (5)$$

Here we term, w_i as the weight of item in an attribute and $|T|$ as the amount of transactions in the database D.

c) *Generating the value of U-gain*

Correspondingly the initial rule R1' is preferred from the rearranged list S and the independent attributes of R1' are derived. By considering the U-factor and the utility value U_i , the value of U-gain for each character attribute is determined.

Definition 3: Item Utility (U_i): In general every character has a precincts of the gain related to that particular character or attribute and is delineated as the Item utility U_i .

Definition 4: Utility table U: A quantity of 'm' utility values U_i are encompassed in the utility table U with the attributes related in the transaction database D. We represent the utility table as:

$$U = \begin{bmatrix} U_1 \\ U_2 \\ \vdots \\ U_m \end{bmatrix} \quad (6)$$

Definition 5: Utility factor (U-factor): The constant value of utility factor (U-factor) is derived by the addition of every utility items (U_i) of the utility table U. We define it as:

$$U - factor = \frac{1}{\sum_{i=1}^m U_i} \quad (7)$$

Consider, m is the amount of attributes involved in the transaction database.

Definition 6: Utility Gain (U-gain): The calculation of an attribute's authentic utility by considering its U-factor is referred as the Utility Gain and we define it as follows:

$$U - gain = U_i \cdot u - factor \quad (8)$$

For every attribute in the association rule R1' the value of U-gain is calculated.

Definition 7: Diminution table: A quantity of 'm' diminution values DM_i are encompassed in the diminution table DM with the attributes related in the transaction database D. We represent the diminution table as:

$$DM = \begin{bmatrix} DM_1 \\ DM_2 \\ \vdots \\ DM_m \end{bmatrix} \quad (9)$$

Definition 8: Diminution factor (D-factor): The constant value of diminution factor (D-factor) is derived by the addition of every diminution items (DM_i) of the diminution table DM. We define it as:

$$D - factor = \frac{1}{\sum_{i=1}^m DM_i} \quad (10)$$

Consider, m is the no of attributes involved in the transaction database.

Definition 9: Diminution Sum (D-sum): The calculation of an attribute's authentic utility by considering its U-factor is referred as the Utility Gain and we define it as follows:

$$D - sum = DM_i \cdot D - factor \quad (11)$$

For every attribute in the association rule R1' the value of D-sum is calculated.

d) *Generation of EUW-score through W-gain, U-gain and D-sum*

From the values derived by calculating W-gain, U-gain and D-sum for the each attribute, they are merged together into one value named as UW-score for every independent association rule.

Definition 10: Exact Utility Weighted Score (EUW-score): EUW-score is derived by computing the proportion between the addition of products of W-gain, U-gain and D-sum for each attribute in the association rule to the total amount of attributes present in the rule.

$$EUW - score = \frac{\sum_{i=1}^{|R|} (W - gain)_i * ((U - gain) - (D - sum))}{|R|} \quad (12)$$

Here, | R | denotes the amount of attributes present in the association rule.

The equations (5),(8) and (11) and (12) aimed to determine the W-gain, U-gain, D-sum and EUW-score. These equations are looped for the remaining association rules R2 ' to R k' involved in the rearranged list S. And for total 'k' number of association rules in the rearranged list S will be calculated with a EUW-Score related to it and the association rules in the rearranged list S are consequently rearranged by taking EUW-score into consideration to get

$$S' = \{R_1'', R_2'', \dots, R_k''\}, \text{ where}$$

$$EUW - Score(R_1'') \geq EUW - Score(R_2'') \geq EUW - Score(R_3'') \dots \geq EUW - Score(R_k'').$$

e) *Deriving the vital association rules by considering EUW-score*

As a final point, we choose certain rules from a group of important weighted utility association rules R_{EUW} in the rearranged list S', whose EUW-Score is greater than the predefined threshold. The

consequential values of the weighted and utility related association rules is given by

$$R_{EUW} = \{R_{EUW1}, R_{EUW2}, R_{EUW3}, \dots, R_{EUWk}\}, \text{ where } k \geq 1 \text{ and } R_{EUW} \subseteq S'.$$

The significant improvement in minimizing number of rules can be observable in following graphs.

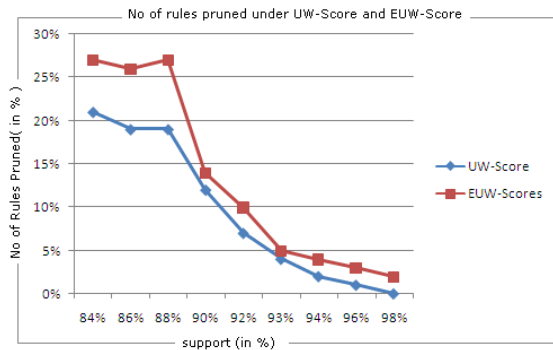


Fig 1 : % of Rules pruned by UW-Score and EUW-Score-line chart representation

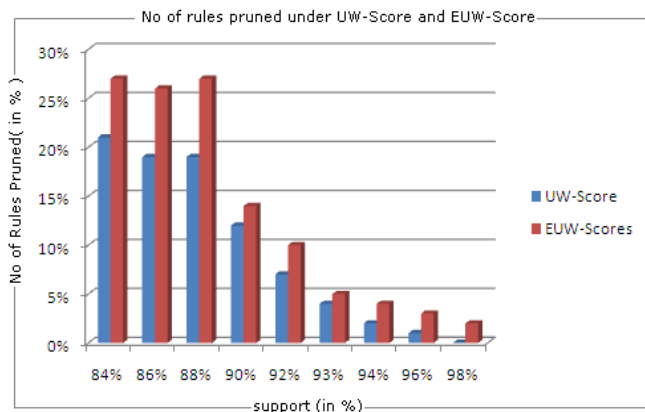


Fig 2 : % of Rules pruned by UW-Score and EUW-Score-bar chart representation

From fig 1 and fig 2, we can observe that number of rules minimized by EUW-Score is significantly better than UW-Score.

IV. CONCLUSION

By considering the weight factor, utility and diminution, the methodology used by us has given a chance to provide a proficient high utility association rules. At the outset, the planned methodology has enabled to make utilization of the conventional Apriori algorithm to create a group of association rules from a database. Depending on weightage (W-gain), utility (U-gain) and diminution (D-sum) complications a joint Exact Utility Weighted Score (EUW-Score) is generated for each association rule extracted. Considering the EUW-Score generated, eventually a subset of notable association rules are derived at.

REFERENCES REFERENCES REFERENCIAS

1. Ali Rajabzadeh Ghatari, Nasibeh Mohamadi, Aida Honarmand, Parviz Ahmadi and Nasibeh Mohamadi, "Recognizing & rioritizing Of Critical Success Factors (CSFs) On Data Mining Algorithm's Implementation In Banking Industry: Evidence From Banking Business System", In Proceedings of EABR & TLC, Prague, Czech Republic, 2009.
2. R. Agrawal, T. Imielinski and A. Swami, "Mining association rules between sets of items in large databases", In proceedings of the international Conference on Management of Data, ACM SIGMOD, pp. 207–216, Washington, DC, May 1993.
3. R. Agrawal and R. Srikant, "Fast algorithms for mining association rules", In Proceedings of 20th International Conference on Very Large Data Bases, Santiago, Chile, pp. 487–499, September 1994.
4. [4] M. S. Chen, J. Han, P.S. Yu, "Data mining: an overview from a database perspective", IEEE Transactions on Knowledge and Data Engineering, vol. 8, no.6, pp. 866–883, 1996.
5. Chun-Jung Chu, Vincent S. Tseng and Tyne Liang, "Mining temporal rare utility Itemsets in large databases using relative utility thresholds", International Journal of Innovative Computing, Information and Control, Vol. 4, no. 11, November 2008.
6. Frawley, W., Piatetsky-Shapiro, G., Matheus, C., "Knowledge Discovery in Databases: An Overview", AI Magazine, fall 1992, pp.213-228, 1992.
7. Guangzhu Yu, Shihuang Shao and Xianhui Zeng, "Mining Long High Utility Itemsets in Transaction Databases", WSEAS Transactions On Information Science & Applications, vol.5, no. 2, pp.202-210, February 2008.
8. A.M.J. Md. Zubair Rahman and P. Balasubram, "Weighted Support Association Rule Mining using Closed Itemset Lattices in Parallel", International Journal of Computer Science and Network security, Vol.9 No. 3 pp. 247-253, March 2009.
9. Hong Yao, Howard J. Hamilton, and Cory J. Butz, "A Foundational Approach to Mining Itemset Utilities from Databases", In Proceedings of the Third SIAM International Conference on Data Mining, pp. 482-486, Orlando, Florida, 2004.
10. Jing Wang, Ying Liu, Lin Zhou, Yong Shi, and Xingquan Zhu, "Pushing Frequency Constraint to Utility Mining Model", In Proceedings of the 7th international conference on Computational Science, pp.685 - 692, Beijing, China, 2007.
11. Jianying Hu and Aleksandra Mojsilovic, "High-utility pattern mining: A method for discovery of high-utility item sets", Pattern Recognition, vol. 40, no. 11, pp.3317-3324, November 2007.
12. Yu-Chiang Li, Jieh-Shan Yeh and Chin-Chen Chang, "Isolated items discarding strategy for discovering high utility itemsets", Data & Knowledge

- Engineering, vol.64, no.1, pp.198-217, January 2008.
13. J. Han and Y. Fu, "Attribute-Oriented Induction in Data Mining," Advances in Knowledge Discovery and Data Mining, AAAI Press/The MIT Press, pp. 399- 421, 1996.
 14. Bing Liu, Wynne Hsu, Ke Wang, and Shu Chen, "Visually Aided Exploration of Interesting Association Rules", In Proceedings of the Third Pacific-Asia Conference on Methodologies for Knowledge Discovery and Data Mining, pp.380 - 389, 1999.
 15. Shankar.S, Babu Nishanth, T. Purusothaman and Jayanthi. S., "A Fast Algorithm for Mining High Utility Itemsets", In proceedings of IEEE International Advance Computing Conference, Patiala,India, 2009.
 16. Ke Sun and Fengshan Bai, "Mining Weighted Association Rules without Preassigned Weights", IEEE Transactions on Knowledge and Data Engineering, Vol. 20, no. 4, April 2008.
 17. Cai. C.H., Fu. A.W.C., Cheng. C. H and Kwong. W.W, "Mining Association Rules with Weighted Items", In Proceedings of the International Symposium on Database Engineering and Applications, pp. 68-77, Cardiff, Wales, UK, July 1998.
 18. Wang. W., Yang. J and Yu. P. S, "Efficient Mining of Weighted Association Rules (WAR)", In Proceedings of the KDD, pp. 270-274, Boston, MA, August 2000.
 19. Lu, S., Hu, H., Li, F, "Mining Weighted Association Rules", Intelligent Data Analysis, vol.5 , no. 3, pp.211 - 225, August 2001.
 20. M. Sulaiman Khan, Maybin Muyebea and Frans Coenen, "Fuzzy Weighted Association Rule Mining with Weighted Support and Confidence Framework", International Workshops on New Frontiers in Applied Data Mining, Osaka, Japan, pp. 49 - 61, May 20-23, 2009.
 21. M. Sulaiman Khan, Maybin Muyebea and Frans Coenen, "A Weighted Utility Framework for Mining Association Rules", In proceedings of European Symposium on Computer Modeling and Simulation, pp.87- 92, Liverpool, September 2008.
 22. Oded Z. Maimon, Lior Rokach, "Decomposition Methodology for Knowledge Discovery and Data Mining: Theory and Applications", World Scientific Publishing Company, ISBN-13: 9789812560797, May 2005.
 23. Feng Tao, Fionn Murtagh and Mohsen Farid, "Weighted Association Rule Mining using Weighted Support and Significance Framework", In Proceedings of the International Conference on Knowledge Discovery and Data Mining, pp.661 - 666, Washington, 2003.
 24. Hong Yaoa and Howard J. Hamilton, "Mining itemset utilities from transaction databases", Data & Knowledge Engineering, vol. 59, no.3, pp. 603-626, December 2006.
 25. Chun-Jung Chua, Vincent S. Tsengb and Tyne Liang, "An efficient algorithm for mining high utility itemsets with negative item values in large databases", Applied Mathematics and Computation, vol. 215, no.2, pp. 767-778, 2009.
 26. Unil Yuna, "Efficient mining of weighted interesting patterns with a strong weight and/or support affinity", Information Sciences, vol.177, no. 17, pp. 3477-3499, September 2007.
 27. Sandhu, P.S.; Dhaliwal, D.S.; Panda, S.N.; Bisht, A.; , "An Improvement in Apriori Algorithm Using Profit and Quantity," Computer and Network Technology (ICCNT), 2010 Second International Conference on , vol., no., pp.3-7, 23-25 April 2010.



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Data mining with Predictive analysis for healthcare sector: An Improved weighted associative classification approach

By Y.Shirisha, S.Siva Shankar Rao, D. Sujatha

Department of CSE

Abstract - Association mining has seen its growth right through data mining during the last few years as it has the ability to search for that entire database that could be of least constraints associated with it. Thus finding such small database sets could be done with the help of predictive analysis method. The paper enlightens the combinational classification of association and classification data mining. For this to happen a new set of constraints need to be introduced namely classification association rule (CAR). Some systems like classification systems with domain experts are the ones that can be associated with. For fields like medicine where a lot many patients consult each doctor, but every patient has got different personal details not necessarily may suffer with same disease. So the doctor may look for a classifier, which could provide all details about every patient and henceforth necessary medications can be provided. However there have been many other classification methods like CMAR, CPAR, MCAR and MMA and CBA. Some advance associative classifiers have also seen growth very recently with small amendments in terms of support and confidence, thereby accuracy. In this paper we proposed a HIT algorithm based automated weight calculation approach for weighted associative classifier.

Keywords : classifier, Association rules, data mining, healthcare, Associative Classifiers, CBA, CMAR, CPAR, MCAR.

GJCST Classification : H.2.8



Strictly as per the compliance and regulations of:



Data mining with Predictive analysis for healthcare sector: An Improved weighted associative classification approach

Y. Shirisha^a, S. Siva Shankar Rao^a, D. Sujatha^b

Abstract - Association mining has seen its growth right through data mining during the last few years as it has the ability to search for that entire database that could be of least constraints associated with it. Thus finding such small database sets could be done with the help of predictive analysis method. The paper enlightens the combinational classification of association and classification data mining. For this to happen a new set of constraints need to be introduced namely classification association rule (CAR). Some systems like classification systems with domain experts are the ones that can be associated with. For fields like medicine where a lot many patients consult each doctor, but every patient has got different personal details not necessarily may suffer with same disease. So the doctor may look for a classifier, which could provide all details about every patient and henceforth necessary medications can be provided. However there have been many other classification methods like CMAR, CPAR, MCAR and MMA and CBA. Some advance associative classifiers have also seen growth very recently with small amendments in terms of support and confidence, thereby accuracy. In this paper we proposed a HIT algorithm based automated weight calculation approach for weighted associative classifier.

Keywords : classifier, Association rules, data mining, healthcare, Associative Classifiers, CBA, CMAR, CPAR, MCAR.

I. INTRODUCTION

A set of steps followed to extract data from the related pattern is termed as data mining. From a haphazard data set it is possible to obtain new data. The predictive modeling approach that simply combines association and classification mining together [3] shows better accuracy [10]. The classification techniques CBA [10], CMAR [9], CPAR [8] out beat the traditional classifiers C4.5, FOIL, RIPPER which are faster but not accurate. Associate classifiers are fit to those model applications which provide support domains in the decisions. However the most suited example for this is medical field where in the data for each patient is required to be stored, with the help of which the system predicts the diseases likely to be affecting the patient. With the system throughput the doctor may decide the medication [6].

Author^a : M.Tech, Department of CSE, ATRI, Parvathapur, Uppal, Hyderabad, India. E-mail : shirisha37@gmail.com

Author^a : Assoc.Prof. Department of CSE, ATRI, Parvathapur, Uppal, Hyderabad, India. E-mail : sivasankarssr@gmail.com

Author^b : Assoc.Prof and HOD Department of CSE, ATRI, Parvathapur, Uppal, Hyderabad, India. E-mail : sujatha.dandu@gmail.com

II. ASSOCIATION CLASSIFICATION

This method of mining is an algorithm based technique with support of some association constraints to retrace the data. The data available is fragmented into 2 parts. One of which is 70% of total data and is taken as training data and the rest taken as another part is used for testing purpose. This technique of data mining is done on data base sets with a record of information.

- Step1: with the help of training set of data produce an association rule set.
- Step2: eliminate all those rules that may cause over fitting.
- Step3: finally we predict the data and check for accuracy and this is said to be the classification phase.

One such example of data base set is:

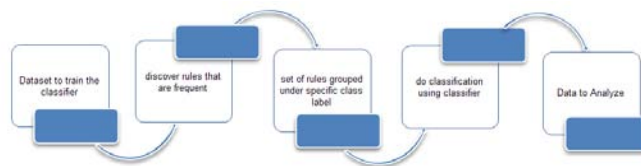


Figure 1 : Associative Classifier for Data Mining

Transaction ID	Items in Transaction
1	ABCDE
2	ACE
3	BD
4	ADE
5	ABCDE

Table 1 : Transactional Database

An association rule is an implication of the form $A \Rightarrow B$, where $A, B \subseteq I$, where I is set of all items, and $A \cap B = \emptyset$. The rule $A \Rightarrow B$ has a support s in the transaction set D if $s\%$ of the transactions in D contain $A \cup B$. The rule $A \Rightarrow B$ holds in the transaction set D with confidence c if $c\%$ of transactions in D that

contain A also contain B. if the threshold point is crossed in terms of confidence then the association rules could be determined. Thus the determined rules form a confidence frame with the help of high strength rules.

On a particular set of domains the AC is performed. A tuple is a collection of m attributes $a_1, a_2, \dots, a_m$ and consists of another special predictive Attribute (Class Label). However these attributes can be categorized depending action methods. Association rule mining is quite different from AC but still undergo following similarities:

- Attribute which is pair (attribute, value), is used in place of Item. For example (BMI, 40) is an attribute in Table 2
- Attribute set is equivalent to Itemset for example ((Age, old), (BP, high)).
- Support count of Attribute (A_i, v_i) is number of rows that matches Attribute in database.
- Support count of Attribute set $(A_i, v_i), \dots, (A_m, v_m)$ is number of rows that match Attribute set in data base.
- An Attribute (A_i, v_i) passes the minsup threshold if support count $(A_i, v_i) \geq \text{min sup}$.

Record ID	Age	Smokes	Hypertension	BMI	Heart Disease
1	42	YES	YES	40	YES
2	62	YES	NO	28	NO
3	55	NO	YES	40	YES
4	62	YES	YES	50	YES
5	45	NO	YES	30	NO

Table 2 : Sample Database for heart patient.

- An Attribute set $((A_i, v_i), \dots, (A_m, v_m))$ passes the threshold if support count $((A_i, v_i), (A_{i+1}, v_{i+1}), \dots, (A_j, v_j)) \geq \text{min sup}$.
- CAR Rules are of form where $c \in ((A_i, v_i), (A_{i+1}, v_{i+1}), \dots, (A_j, v_j)) \rightarrow c$ Class-Label. Where Left hand side is itemset and right hand side is class. And set of all attribute and class label together ie $((A_i, v_i), \dots, (A_j, v_j), c)$ is called rule attribute.
- Support count of rule attribute $((A_i, v_i), (A_{i+1}, v_{i+1}), \dots, (A_j, v_j), c)$ is number of rows that matches item in database. Rule attribute $((A_i, v_i), (A_{i+1}, v_{i+1}), \dots, (A_j, v_j), c)$ passes the threshold if support count of $((A_i, v_i), (A_{i+1}, v_{i+1}), \dots, (A_j, v_j), c) \geq \text{min sup}$.

An important subset of rules called class association rules (CARs) are used for classification

purpose since their right hand side is used for attributes. Its simplicity and accuracy makes it efficient and friendly for end user. Whenever any amendments need to be done in a tree they can be made without affecting the other attributes.

III. ADVANCEMENTS IN CAR RULE GENERATION

The accuracy of the classification however depends on the rules implied in the classification. To overcome CARs rules inaccuracy in some cases, a new advanced ARM in association with classifiers has been developed. This new advanced technique provides high accuracy and also improves prediction capabilities.

a) An Associative Classifier Based On Positive And Negative Approach

Negative association rule mining and associative classifiers are two relatively new domains of research, as the new associative amplifiers that take advantage of it. The positive association rule of the form $X \rightarrow Y$ to $\neg X \rightarrow Y$, $X \rightarrow \neg Y$ and $\neg X \rightarrow \neg Y$ with the meaning X is for presence and $\neg X$ is for absence. Based on correlation analysis the algorithm uses support confidence Instead of using support-confidence framework in the association rule generation. Correlation coefficient measure is added to support confidence framework as it measures the strength of linear relationship between a pair of two variables. For two

variables X and Y it is given by $\rho = \frac{\text{con}(X, Y)}{\sigma_X \sigma_Y}$, where

$\text{con}(X, Y)$ represents the covariance of two variables and σ_X stand for standard deviation.

The range of values for ρ is between -1 to +1, when it is +1 the variables are perfectly correlated, if it is -1 the variables are perfectly independent then equals to 0. when positive and negative rules are used for classification in UCI data sets encouraging results will obtain. Negative association rules are effective to extract hidden knowledge. And if they are only used for classification, accuracy decreases.

b) Temporal associative classifiers

As data is not always static in nature, it changes with time, so adopting temporal dimension to this will give more realistic approach and yields much better results as the purpose is to provide the pattern or relationship among the items in time domain. For example rather than the basic association rule of $\{\text{bread}\} \rightarrow \{\text{butter}\}$ mining from the temporal data we can get that the support of $\{\text{bread}\} \rightarrow \{\text{butter}\}$ raises to 50% during 7 pm to 10 pm everyday [3], as These rules are more informative they are used to make a strategic decision making. Time is an important aspect in temporal database.

1. Scientific, medical, dynamic systems, computer network traffic, web logs, markets, sales, transactions, machine/device performance, weather/climate, telephone calls are examples of time ordered data.
2. The volume of dynamic time-tagged data is therefore growing, and continuing to grow,
3. as the monitoring and tracking of real-world events frequently require repeated measurements.3.Data mining methods require some modification to handle special temporal relationships ("before", "after", "during", "in summer", "whenever X happens").
4. Time-ordered data lend themselves to prediction – what is the likelihood of an event, given the preceding history of events? (e.g., hurricane tracking, disease epidemics)
5. Time-ordered data often link certain events to specific patterns of temporal behavior (e.g., network intrusion breaks INS).The new type of AC called Temporal Associative Classifier is being proposed in [3] to deal with above such situation. CBA, CMAR and CPAR are modified with temporal dimension and proposed TCBA, TCMAR and TCPAR. To compare the classifying accuracy and execution time of the three algorithms using temporally modified data set of UCI machine learning data sets an experiment has performed and conclusions are:

- i. TCPAR performs better than TCMAR and TCBA as it is time consuming for smaller support values but improves in run- time performance as the support increases.
- ii. Using data set the accuracy is calculated for each algorithm. The average accuracy of TCPAR is found little better than TCMAR.
- iii. The temporal counterpart of all the three associative classifiers has shown improved classification accuracy as compare to the non-temporal associative classifier. Time-ordered data lend themselves to prediction like what is the likelihood of an event e.g., (hurricane tracking, disease epidemics). The temporal data is useful in predicting the disease in different age group.

c) Associative Classifier Using Fuzzy

Association Rule: The quantitative attributes are one of preprocessing step in classification. for the data which is associated with quantitative domains such as income, age, price, etc., in order to apply the Apriori-type method association rule mining needs to partition the domains. Thus, a discovered rule $X \rightarrow Y$ reflects association between interval values of data items. Examples of such rules are "Fruit [1-5kg] $\rightarrow$ Meat [5-20\$]", "Income [20-50k\$] $\rightarrow$ Age [20-30]", and so on [ZC08]. As the record belongs to only one of the set results in sharp boundary problem which gives rise to the notion of fuzzy association rules (FAR).The

semantics of a fuzzy association rule is richer and natural language nature, which are deemed desirable. For example, "low-quantity Fruit $\rightarrow$ normal-consumption Meat" and "medium Income $\rightarrow$ young Age" are fuzzy association rules, where X's and Y's are fuzzy sets with linguistic terms (i.e., low, normal, medium, and young).An associative classification based on fuzzy association rules (namely CFAR) is proposed to overcome the "sharp boundary" problem for quantitative domains. Fuzzy rules are found to be useful for prediction modeling system in medical domain as most of the attributes are quantitative in nature hence fuzzy logic is used to deal with sharp boundary problems.

i. Defining Support and confidence measure

New formulae of support and confidence for fuzzy classification rule $F \rightarrow C$ are as follows:

$$\text{support}(F \rightarrow C) = \frac{\text{Sum of membership values of antecedent with class label C}}{\text{Total No. of Records in the Database}}$$

$$\text{confidence}(F \rightarrow C) = \frac{\text{Sum of membership values of antecedent with class label C}}{\text{Sum of membership values of antecedent for all class label}}$$

d) Weighted Associative Classifiers

Based on the different features weights are allotted based on this classifier. Every attribute varies in terms of importance.it also important to know that with the capabilities of predicting the weights may be altered. A weighted associative classifiers consists of training dataset $T=\{r1,r2, r3.... ri....\}$ with set of weight associated with each {attribute, attribute value} pair. Each with recordri is a set of attribute value and a weight w_i attached to each attribute of rituple / record. Aweighted framework has record as a triple $\{a_i, v_i, w_i\}$ where attribute a_i is having value v_i and weight w_i , $0 < w_j \leq 1$. Thus with the help of weights one can easily determine its predicting ability. With this weighted rules like "medium Income young Age", " $\{(Age, >62), (BMI, <45), (Boold_pressur, <95-135)\}$ ", Heart Disease, (Income[20,000-30,000]Age[20-30]) could become the criteria of determination. Weights of data as per table 2 are recorded in table 4 with the help of weights mentioned in the table 5 for predicting attributes.

Record ID	Age	Smokes	Hypertension	BMI	Record Weight
1	42	YES	YES	40	0.6
2	62	YES	NO	28	0.42
3	55	NO	YES	40	0.52
4	62	YES	YES	50	0.67
5	45	NO	YES	30	0.45

Table 4 : Relational Database with record weight

Here we measure the record weight using following

$$recordWeight = \frac{\sum_{i=1}^{|I|} w_i}{|I|}$$

Here

$|I|$ is total number of items in record

i is an item in record

w_i is weight of the item i

i. *Measuring weights using HITS algorithm*

a. *Ranking Transactions with HITS*

A database of transactions can be depicted as a bipartite graph without loss of information. Let $D = \{T_1, T_2, \dots, T_m\}$ be a list of transactions and $I = \{i_1, i_2, \dots, i_n\}$ be the corresponding set of items. Then, clearly D is equivalent to the bipartite graph $G = (D, I, E)$ where

$$E = \{(T, i) : i \in T, T \in D, i \in I\}$$

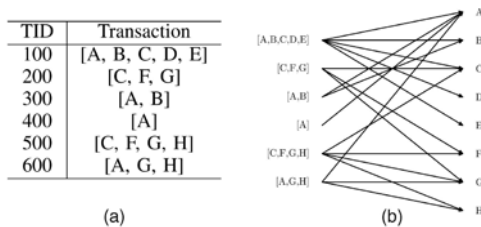


Fig1 : The bipartite graph representation of a database
(a) Database (b) Bipartite graph

Example 1: Consider the database shown in Fig. 1a. It can be equivalently represented as a bipartite graph, as shown in Fig. 1b. The graph representation of the transaction database is inspiring. It gives us the idea of applying link-based ranking models to the evaluation of transactions. In this bipartite graph, the support of an item i is proportional to its degree, which shows again that the classical support does not consider the difference between transactions. However, it is crucial to have different weights for different transactions in order to reflect their different importance. The evaluation of item sets should be derived from these weights. Here comes the question of how to acquire weights in a database with only binary attributes. Intuitively, a good transaction, which is highly weighted, should contain many good items; at the same time, a good item should be contained by many good transactions. The reinforcing relationship of transactions and items is just like the relationship between hubs and authorities in the HITS model [3]. Regarding the transactions as “pure” hubs and the items as “pure” authorities, we can apply HITS to this bipartite graph. The following equations are used in iterations:

$$auth(i) = \sum_{T:i \in T} hub(T), \quad hub(T) = \sum_{i:i \in T} auth(i) \dots (1)$$

When the HITS model eventually converges, the hub weights of all transactions are obtained. These weights represent the potential of transactions to contain high-value items. A transaction with few items may still be a good hub if all component items are top ranked. Conversely, a transaction with many ordinary items may have a low hub weight.

b. *W-support - A New Measurement*

Item set evaluation by support in classical association rule mining [1] is based on counting. In this section, we will introduce a link-based measure called w-support and formulate association rule mining in terms of this new concept.

The previous section has demonstrated the application of the HITS algorithm [3] to the ranking of the transactions. As the iteration converges, the authority weight $auth(i) = \sum_{T:i \in T} hub(T)$ represents the

“significance” of an item i , accordingly, we generalize the formula of $auth(i)$ to depict the significance of an arbitrary item set, as the following definition shows:

Definition 1: The w-support of an item set X is defined as

$$wsup p(X) = \frac{\sum_{T:X \subseteq T \wedge T \in D} hub(T)}{\sum_{T:T \in D} hub(T)} \dots (2)$$

Where $hub(T)$ is the hub weight of transaction T . An item set is said to be significant if its w-support is larger than a user specified value. Observe that replacing all $hub(T)$ with 1 on the right hand side of (2) gives $supp(X)$. Therefore, w-support can be regarded as a generalization of support, which takes the weights of transactions into account. These weights are not determined by assigning values to items but the global link structure of the database. This is why we call w-support link based. Moreover, we claim that w-support is more reasonable than counting-based measurement.

ID	Symptoms	Weight
1	Age<40	0.437
2	40<Age<58	0.375
3	Age>58	0.185
4	Smokes=yes	0.68
5	Smokes=no	0.31
6	Hypertension=yes	0.5
7	Hypertension=no	0.5
8	BMI<=25	0.23
9	26<=BMI<=30	0.2
10	31<=BMI<=40	0.43
11	BMI>40	0.125

Table 5 : weights measured using proposed algorithm

ii. *Defining Support and confidence measure*

New formulae of support and confidence for classification rule $X \rightarrow \text{class_Label}$, where X is set of weighted items, is as follows: Weighted Support: Weighted support WSP of rule $X \rightarrow \text{class_Label}$, where X is set of non empty subsets of attribute value set, is fraction of weight of the record that contain above attribute-value set relative to the weight of all transactions.

$$v1 = \sum_{i=1}^{|X|} \text{weight}(r_i)$$

$$v2 = \sum_{k=1}^{|n|} \text{weight}(r_i)$$

$$\text{WSP}(X \rightarrow \text{Class_Label}) = \frac{v1}{v2}$$

The classification accuracy improvement for Weighted associative classifier with proposed weight measurement approach can be observable in fig 2 and fig 3.

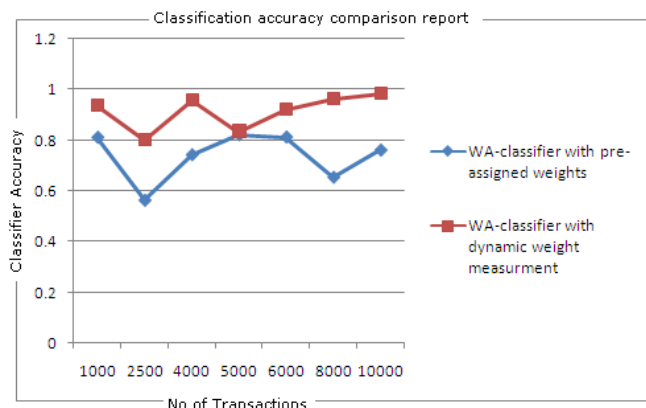


Fig 2 : A line chart representation of classification accuracy differences between Traditional Weighted associative classifier and proposed weighted associative classifier

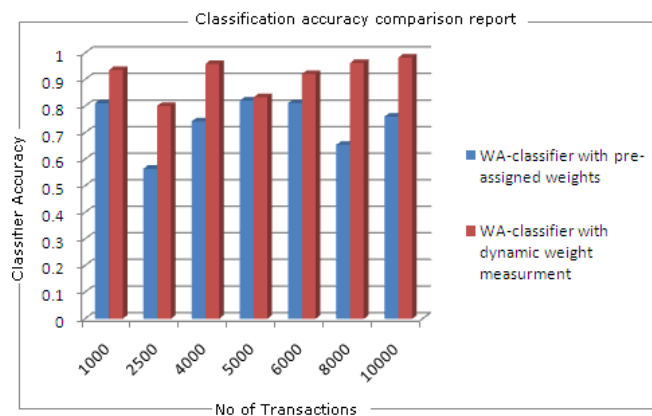


Fig 3 : A bar chart representation of classification accuracy differences between Traditional Weighted associative classifier and proposed weighted associative classifier

IV. REFINING SUPPORT AND CONFIDENCE MEASURES TO VALIDATE DOWNWARD CLOSURE PROPERTY

The downward closure property is the key part of Apriori algorithm. It states that any super set can't be frequent unless and until its itemset isn't frequent. The itemsets that are already found to be frequent are added with new items based on the algorithm. However changes in support and confidence shall not show its effect on this property and also AC associated with advanced rule developer. The terms support and confidence are to be replaced with weighted support and weighted confidence respectively in WAC which elicits that weighted support helps maintain weighted closure property.

V. CONCLUSION

This advanced AC method could be applied in real time scenario to get more accurate results. This needs lot of prediction to be done based on its capabilities which could be improved. It finds its major application in the field of medical where every data has an associated weight. The proposed HIT algorithm based weight measurement model is significantly improving the quality of classifier.

REFERENCES

1. Sunita Soni, Jyothi Pillai, O.P. Vyas An Associative Classifier Using Weighted Association Rule 2009 International Symposium on Innovations in natural Computing, World Congress on Nature & Biologically.
2. Zuoliang Chen, Guoqing Chen BUILDING AN ASSOCIATIVE CLASSIFIER BASED ON FUZZY ASSOCIATION RULES International Journal of Computational Intelligence Systems, Vol.1, No. 3 (August, 2008), 262 – 273.

3. RanjanaVyas, Lokesh Kumar Sharma, Om Prakashvyas, Simon ScheiderAssociative Classifiers for Predictive analytics: Comparative Performance Study, second UKSIM European Symposium on Computer Modeling and Simulation 2008.
4. E.RamarajN.VenkatesanPositive and Negative Association Rule Analysis in Health Care Database ,IJCSNS International Journal of Computer Science and Network Security, VOL.8 No.10, October 2008,325-330.
5. Khan, M.S. Muyebe, M. Coenen, F A Weighted Utility Framework for Mining Association Rules, Symposium Computer Modeling and Simulation, 2008. EMS '08. Second UKSIM European, page(s): 87-92.
6. FadiThabtah, A review of associative classification mining, The Knowledge Engineering Review, Volume 22, Issue 1 (March 2007), Pages 37-65, 2007.
7. LuizaAntonie, University of Alberta, Advancing Associative Classifiers - Challenges and Solutions, Workshop on Machine Learning, Theory, Applications, Experiences 2007
8. Feng Tao, FionnMurtagh and Mohsen Farid. Weighted Association Rule Mining using Weighted Support and Significance Framework Proceedings of the ninth ACM SIGKDD international conference on Knowledge discovery and data mining 2003, Pages:661-666 Year of Publication: 2003.
9. Yin, X. & Han, J. CPAR: Classification based on predictive association rule. In Proceedings of the SIAM International Conference on Data Mining. San Francisco, CA: SIAM Press, 2003,pp. 369-376.
10. W. Li, J. Han, and J. Pei. CMAR: Accurate and efficient classification based on multiple class association rules. In ICDM'01, pp. 369-376, San Jose, CA, Nov.2001.
11. Liu,B Hsu. W. Ma, Integrating Clasification and association rule mining .Proceeding of the KDD, 1998(CBA) pp 80-86.
12. Cláudia M. Antunes, and Arlindo L. Oliveira Temporal Data Mining: an overview.
13. Antonie, M. &Zaiane, O. An associative classifier based on positive and negative rules. In Proceedings of the 9th ACM SIGMOD Workshop on Research Issues in Data Mining and Knowledge Discovery, 2004,pp 64-69.



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Fuzzy and Swarm Intelligence for Software Cost Estimation

By Srinivasa Rao.T, Prasad Reddy P.V.G.D, Hari Ch.V.M.K

GITAM University

Abstract - Software cost estimation is the process of predicting the amount of time, effort and resources required to complete the project successfully. Software development is a collection of activities includes feasibility study, analysis, design, coding, testing, implementation and maintenance. Each phase requires resources- people, time, software and hardware which should be predicted well before the software development. The prediction means lot of uncertainty. So far many models are proposed by using Fuzzy Logic, Neural Networks, Machine Learning, Regression analysis and Soft Computing techniques. In this paper we are proposed a new model structure basing on Alaa F. Sheta using Fuzzy logic for controlling prediction uncertainty and the parameters of the cost model tuned by using swarm intelligence-Particle Swarm Optimization. The proposed model results are verified with NASA software dataset and results are compared with the existing models. The Results show that the value of MARE (Mean Absolute Relative Error) applying fuzzy-swarm intelligence was substantially lower than MARE of other models exists in the literature.

Keywords : *Software Cost Estimation, Swarm Intelligence, Fuzzy Logic, COCOMO, Particle Swarm Optimization.*

GJCST Classification : *D.2.9*



FUZZY AND SWARM INTELLIGENCE FOR SOFTWARE COST ESTIMATION

Strictly as per the compliance and regulations of:



Fuzzy and Swarm Intelligence for Software Cost Estimation

Srinivasa Rao.T^α, Prasad Reddy P.V.G.D^α, Hari Ch.V.M.K^β

Abstract - Software cost estimation is the process of predicting the amount of time, effort and resources required to complete the project successfully. Software development is a collection of activities includes feasibility study, analysis, design, coding, testing, implementation and maintenance. Each phase requires resources- people, time, software and hardware which should be predicted well before the software development. The prediction means lot of uncertainty. So far many models are proposed by using Fuzzy Logic, Neural Networks, Machine Learning, Regression analysis and Soft Computing techniques. In this paper we are proposed a new model structure basing on Alaa F. Sheta using Fuzzy logic for controlling prediction uncertainty and the parameters of the cost model tuned by using swarm intelligence-Particle Swarm Optimization. The proposed model results are verified with NASA software dataset and results are compared with the existing models. The Results show that the value of MARE (Mean Absolute Relative Error) applying fuzzy-swarm intelligence was substantially lower than MARE of other models exists in the literature.

Keywords : Software Cost Estimation, Swarm Intelligence, Fuzzy Logic, COCOMO, Particle Swarm Optimization.

I. INTRODUCTION

Software project management is collection of two activities: Project Planning and Project Monitoring and control. Planning is predicting the activities that must be done before starting development work. Once project work is started it is the responsibility of project manager to monitor the work and see the goal-high quality of software must be produced with low cost and within a time and budget. The input for the planning is SRS-Software Requirement Specification Document and output is project plan mainly includes Cost estimation and Schedule estimation.

Software cost estimation is the process of predicting the amount of time required to build a software system. The time is measured in terms of Person-Months (PM's) which is later on converted into dollar cost. The basic input for the cost model is size measured in terms of KDLOC (Kilo Delivered Lines Of Code) and set of Cost parameters. The advantage of cost estimation is Cost benefit analysis, proper resource utilization (software, hardware and people), staffing plans, functionality trade-offs, risks and modify budget.

*Author ^{αβ} : GITAM Institute Of Technology, GITAM University.
E-mail : tsr.etl@gmail.com,*

*Author ^α : Department Of CS&SE, Andhra University.
E-mail : kurmahari@gmail.com*

The software cost estimation problem deserves a special attention because of development of product is unique under taking results in uncertainty, with increased size of software projects estimation mistakes could cost lot in terms of resources allocated to the project.

II. BACKGROUND

In this section we briefly discuss the COCOMO (Constructive Cost Model), Fuzzy Logic and Swarm Intelligence-Particle Swarm Intelligence.

a) COCOMO

[Boehm, 1981][4] described COCOMO as a collection of three variants, they are Basic model, Intermediate model, and Detailed model. Boehm described three development modes and Organic is for relatively simple projects, Semidetached is for relatively intermediate projects, Embedded for a project developed under tight constraints.

The Basic COCOMO Model computes effort E as function of program size, and it is same as single variable method. The Effort calculated using the following equation

$$\text{Effort} = a * (\text{size})^b \quad (1)$$

Where a and b are the set of values depending on the complexity of software (for organic projects a=2.4, b=1.05, for semi-detached a=3.0, b=1.1.2 and for embedded a=3.6, b=1.2).

An Intermediate COCOMO model effort is E is function of program size and set of cost drivers or effort multipliers. The Effort calculated using the following equation

$$\text{Effort} = a * (\text{size})^b * \text{EAF} \quad (2)$$

where a and b are the set of values depending on the complexity of software (for organic projects a=3.2, b=1.05, for semi-detached a=3.0, b=1.1.2 and for embedded a=2.8, b=1.2) and EAF (Effort Adjustment Factor) which is calculated using 15 cost drivers. Each cost driver is rated from ordinal scale ranging from low to high.

In Detailed COCOMO the effort E is function of program size and a set of cost drivers given according to each phase of software life cycle. The phases used in detailed COCOMO are requirements planning and product design, detailed design, code and unit test, and

integration testing. The weights defined accordingly. The Effort calculated using the following equation

$$\text{Effort} = a * (\text{size})^b * \text{EAF} * \sum(W_i) \quad (3)$$

Boehm and his colleagues have refined and updated COCOMO called as COCOMO II. It is a collection of three variants, Application composition model, early design model, and Post architecture model.

b) Fuzzy Logic

A fuzzy set is a set with a smooth boundary. Fuzzy set theory generalizes classical set theory to allow partial membership[5,6]. The best way to introduce fuzzy sets is to start with a limitation of classical sets. A set in classical set theory always has a sharp boundary because membership in a set is a black-and-white concept, i.e. an object either completely belongs to the set or does not belongs to the set at all. The degree of membership in a set is expressed by a number between 0 and 1; 0 means entirely not in the set, 1 means completely in the set, and a number in between means partially in the set. This way a smooth and gradual transition from the region outside the set to those in the set can be described. A fuzzy set is thus defined by a function that maps objects in a domain of concern to their membership value in the set. Such a function is called the Membership Function and usually denoted by the Greek symbol μ . The membership function of a fuzzy set A is denoted by μ_A , and the membership value of x in A is denoted by $\mu_A(x)$. The domain of membership function, which is the domain of concern from which elements of the set are drawn, is called the Universe Of Discourse. We may identify meaningful lower and upper bounds of the membership functions. Membership functions of this type are known as interval values fuzzy sets. The intervals of the membership functions are also fuzzy then it is known as interval Type-2 fuzzy sets.

c) Swarm Intelligence-Particle Swarm Optimization

Swarm Intelligence (SI) is an innovative distributed intelligent paradigm for solving optimization problems that originally took its inspiration from the biological examples by swarming, flocking and herding phenomena in vertebrates. Particle Swarm Optimization (PSO) incorporates swarming behaviors observed in flocks of birds, schools of fish, or swarms of bees, and even human social behavior, from which the idea is emerged. PSO is a population-based optimization tool, which could be implemented and applied easily to solve various function optimization problems, or the problems that can be transformed to function optimization problems. Particle Swarm Optimization was first introduced by Dr. Russell C. Eberhart and Dr. James Kennedy in 1995. As described by Eberhart and Kennedy, the PSO algorithm is an adaptive algorithm based on a social-psychological metaphor; a population of individuals (referred to as particles) adapts by

returning stochastically toward previously successful regions. The basic concept of PSO lies in accelerating each particle towards its Pbest and Gbest locations with a random weighted acceleration at each time. The modification mbns of the particles positions can be mathematically modeled according to the following equations:

$$V_i^{k+1} = w * V_i^k + c_1 * \text{rand}() * (V_{pbest} - S_i^k) + c_2 * \text{rand}() * (V_{gbest} - S_i^k) \quad (4)$$

$$S_i^{k+1} = S_i^k + V_i^{k+1} \quad (5)$$

Where, S_i^k is current search point, S_i^{k+1} is modified search point, V_i^k is the current velocity, V_i^{k+1} is the modified velocity, V_{pbest} is the velocity based on Pbest, V_{gbest} = velocity based on Gbest, w is the inertia weight, c_i is the weighting factors, $\text{rand}()$ are uniformly distributed random numbers between 0 and 1. In order to guide the particle effectively in the search space, the maximum moving distance during each iteration must be changed in between the maximum velocity $[-V_{max}, V_{max}]$.

III. LITERATURE REVIEW

In this section we discuss the some previous models proposed using Genetic Algorithms[8], Fuzzy models[9], Soft-Computing Techniques[10], Computational Intelligence Techniques[3], Heuristic Algorithms, Neural Networks[7], Radial Basis[11], and Regression[1,2,4].

The Cost of product is a function of many parameters which are Size (coding size), Cost Drivers and Methodology used in the project. The Walston Felix uses 36 cost drivers, 16 by Boheam and 30 other factors considered by the Bailey-Basili for the cost estimation. The parameters are estimated by using regression analysis and the effort equation is[1]

$$E = 5.5 + 0.73(KLOC)^{1.16} \quad (6)$$

Where E is effort and KLOC is kilo lines of code-coding size

The Alaa F. Sheta proposed two new model structures by using genetic algorithms for tuning parameters. The Model 1 and 2 equations is

$$\text{Effort} = 3.1938(DLOC)0.8209 - 0.1918(ME) \text{ for model 1} \quad (7)$$

$$\text{Effort} = 3.3602(DLOC)0.8116 - 0.4524(ME) + 17.8025 \text{ for model 2} \quad (8)$$

Where ME is the methodology used in the project.

The Harish proposed two model structures based on triangular fuzzy sets [13]. Interval Type-2 fuzzy logic, Particle Swarm Optimization for is proposed by Prasad Reddy. [12].

IV. PROPOSED METHODOLOGY AND ALGORITHM

a) Methodology

The uncertainty about cost estimation is usually quite high, because of prediction of basic element size, cost drivers and other parameters. By introducing some modifications in the interval type-2 fuzzy logic we can control the uncertainty. In the present work fuzzy sets are used for modeling uncertainty and imprecision in an efficient way. The inputs of the standard cost model include an estimation of project size and evaluation of the parameters, rather than a single number, the software size can be regarded as a fuzzy set yielding the cost estimate also in the form of a fuzzy set. We emphasize a way of propagation of uncertainty and ensuring violation of the resulting effort. Fuzzy sets create a more flexible, high versatile development environment. They generate a feedback also the resulting uncertainty of the results. The decision-maker is no longer left with a single variable estimate which could be highly misleading in many cases and lead to the belief as a to the relevance of the obtained results.

In the present work on the proposed models the parameter tuning is done by using Particle Swarm Optimization. For each particle position with values of tuning parameters, fitness function is evaluated with an objective to minimize the fitness function. The objective is to minimize or maximize fitness function. The particles moving towards optimal parameters by doing several iterations until particles exhaust or derivative of velocity becomes nearly zero then we get the optimal parameters which are later used for effort estimation.

b) Proposed Model

The Proposed model consists of three major components. First component is fuzzification process which identifies the suitable firing intervals for the input parameter size. Second component is parameter tuning using particle swarm optimization. Finally effort estimation done through weighted average defuzzification method using the results obtained in first and second steps.

i. Fuzzification process

The input size is fuzzified by using two triangular fuzzy sets. The Triangular member function is shown below.

$$\mu_p(\text{size}) = \begin{cases} 0 & \text{size} \leq L \\ \frac{\text{size} + L}{2L} & -L \leq \text{size} \leq L \\ 1 & \text{size} > L \end{cases} \quad (9)$$

Where L is the Mean of input sizes

After fuzzification of the data set find the shaded regions (overlap) of the left hand side and right hand side those are called meaningful lower and upper bounds of the data set- Foot Print of Uncertainty. The

means of the Foot Print of Uncertainty as firing intervals.

ii. Parameter tuning using Particle Swarm Optimization

The effort equation we considered is Alaa F. Sheta model-2

$$\text{Effort} = a * (\text{Size})^b + c * (\text{ME}) + d \quad (10)$$

The parameters a, b, and c of the equation 10 are tuned by using particle swarm optimization with inertia weight and MARE as the fitness function (minimize).

iii. Defuzzification

The defuzzification is done through weighted average method is as shown below

$$E = \{w_1 * [(a * \alpha^b) + c * (\text{ME}) + d] + w_2 * [(a * m^b) + c * (\text{ME}) + d] + w_3 * [(a * \beta^b) + c * (\text{ME}) + d]\} / w_1 + w_2 + w_3 \quad (11)$$

Where w_i is the weighting factor and α , m , and β are the fuzzified sizes obtained from triangular member function.

V. RESULTS AND DISCUSSIONS

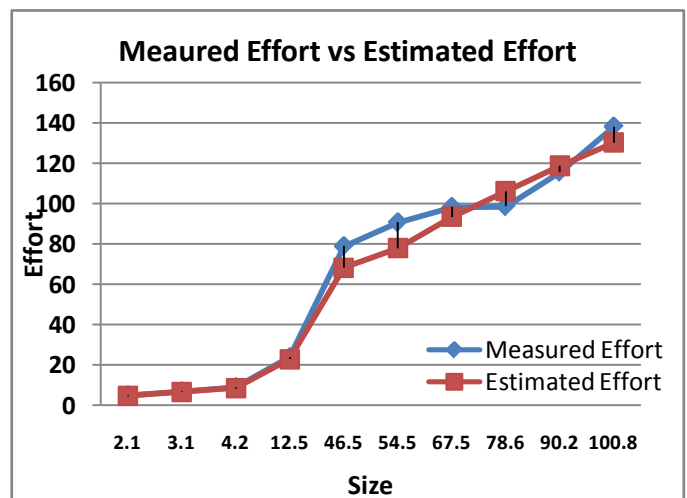
NASA dataset is considered for experimentation. The firing intervals obtained after the fuzzification are [0.7362, 0.8998]. The parameters obtained after tuning PSO methodology $a=3.131606$, $b=0.820175$, $c=0.045208$ and $d=-2.020790$. While performing defuzzification $w_1=1, w_2=10$ and $w_3=10$. The following table 1 shows the efforts the proposed model. The estimated efforts are very close to the measured efforts.

The proposed model results are compared with the existing models in the literature and the results are shown in the following table 2.

The performance measure considered here is Mean Absolute Relative Error (MARE)

$$\% \text{ MARE} = \text{mean}$$

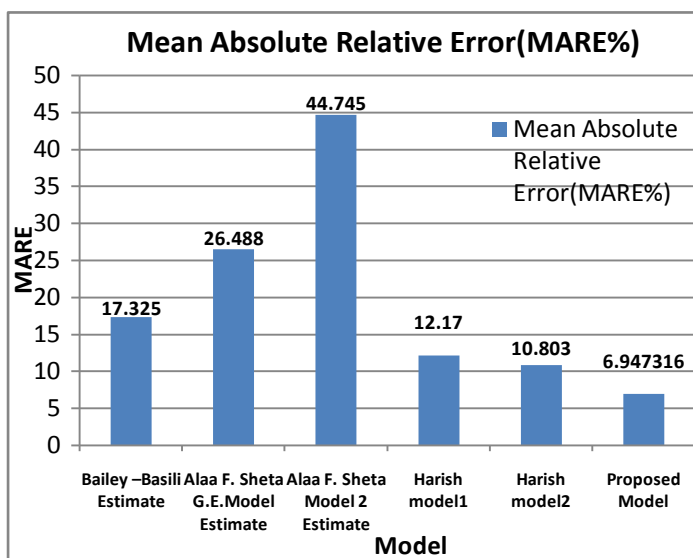
$$\left[\frac{\text{abs}(\text{Measured Effort} - \text{Estimated Effort})}{(\text{measured effort})} \right] \times 100$$



The MARE of various models is shown below.

Model	Mean Absolute Relative Error(MARE%)
Bailey –Basili Estimate	17.325
Alaa F. Sheta G.E.Model Estimate	26.488
Alaa F. Sheta Model 2 Estimate	44.745
Harish model1	12.17
Harish model2	10.803
Proposed Model	6.947316

The Results show that the value of MARE (Mean Absolute Relative Error) applying fuzzy-swarm intelligence was substantially lower than MARE of other models exists in the literature.



VI. CONCLUSION

Software cost estimation is based on a probabilistic model and hence it does not generate exact values. However availability of good historical data coupled with a systematic technique can generate better results. In this paper we proposed new model structure to estimate the software cost (Effort) estimation. Fuzzy sets is used for modeling uncertainty and impression to better the effort estimation and particle swarm optimization for tuning parameters. It is observed from the results that Fuzzy-Swarm intelligence gives accurate results when juxtaposed with its other counterparts. On testing the performance of the model in terms of the MARE the results were found to be useful.

REFERENCES REFERENCES REFERENCIAS

1. John w. Bailey and Victor R.Basili,(1981) "A meta model for software development resource expenditures", proceedings of the Fifth International Conference on Software Engineering, CH-1627-9/81/0000/0107500.75@IEEE, pp:107-129,1981.C

2. Chris F. Kemerer, (1987) "An Empirical Validation of Software Cost Estimation Models", Management Of Computing-Communications of ACM, Vol:30 No. 5, pp:416-429, May 1987.
3. KrishnaMurthy.S and Douglas Fisha (1995), "Machine Learning Approaches to estimating software development effort", IEEE Transactions On Software Engineering, Vol. 21, No. 2, pp: 126-137, February 1995.
4. Barry Boehm, Chris Abts, Sunita Chulani(1998), Software Development Cost Estimation Approaches – A Survey1,IBM, pp: 1-46 ,1998.
5. Hao Ying,(1998), The Takagi-Sugeno fuzzy controllers using the simplified linear control rules are nonlinear variable gain controllers, ELSEVIER: Automatica, Vol:34,No.2,pp.157-167, 1998.
6. Magne Jørgensen,(2005), Evidence-Based Guidelines for Assessment of Software Development Cost Uncertainty, IEEE Transactions On Software Engineering, Vol. 31, No. 11, pp: 942-954, November 2005.
7. Xishi Huang, Danny Ho, Jing Ren, Luiz F. Capretz (2005), Improving the COCOMO model using a neuro-fuzzy approach, DOI:10.1016 /j.asoc. 2005.06.007, Applied Soft Computing 7 (2007) pp: 29–40, @2005 Elsevier.
8. Alaa F. Sheta, (2006), Estimation of the COCOMO Model Parameters Using Genetic Algorithms for NASA Software Projects, Journal of Computer Science 2, (2) pp:118-123, 2006.
9. Alaa Sheta (2006),Software Effort Estimation and Stock Market Prediction Using Takagi-Sugeno Fuzzy Models, Proceedings of 2006 IEEE International Conference on Fuzzy Systems, 0-7803-9489-5/06/IEEE 2006,pp:171-178,July 16-21, 2006.
10. Alaa Sheta, David Rine and Aladdin Ayesh (2008), Development of Software Effort and Schedule Estimation Models Using Soft Computing Techniques, IEEE Transaction, 978-1-4244-1823-7/08/pp: 1283-1289@2008 IEEE.
11. Prasad Reddy P.V.G.D, Sudha, K.R., Rama Sree .P, Ramesh .S.N.S.V.S.C. (2010), Fuzzy Based Approach for Predicting Software Development Effort, International Journal of Software Engineering (IJSE), Vol.1 No.1,pp:1-11,2010.
12. Prasad Reddy P.V.G.D (2010), Particle Swarm Optimization In The Fine Tuning Of Fuzzy Software Cost Estimation Models, International Journal of Software Engineering (IJSE), Vol.1 No.2, pp: 12-23,2010.
13. Harish Mittal and Pradeep Bhatia, (2007), Optimization Criteria for Effort Estimation using Fuzzy Technique, CLEI Electronic Journal, Vol. 10, Number 1, Paper 2,pp: 1-11, June 2007.

Table 1 : Estimated effort of the proposed model.

Size (m)	Methodology	Measured Effort	$\alpha = 0.7362m$	m	$\beta = 0.8998m$	$a*\alpha^b + c*ME + d$	$a*m^b + c*ME + d$	$a*\beta^b + c*ME + d$	Estimated Effort
2.1	28	5	1.546	2.100	1.890	3.722	5.000	4.523	4.712
3.1	26	7	2.282	3.100	2.789	5.316	7.075	6.418	6.679
4.2	19	9	3.092	4.200	3.779	6.742	8.999	8.156	8.490
12.5	27	23.9	9.203	12.500	11.248	18.535	24.055	21.994	22.811
46.5	19	79	34.233	46.500	41.841	55.630	71.846	65.790	68.190
54.5	20	90.8	40.123	54.500	49.039	63.572	82.044	75.145	77.879
67.5	29	98.4	49.694	67.500	60.737	76.386	98.400	90.179	93.437
78.6	35	98.7	57.865	78.600	70.724	86.911	111.853	102.538	106.229
90.2	30	115.8	66.405	90.200	81.162	97.125	125.048	114.620	118.753
100.8	34	138.3	74.209	100.800	90.700	106.636	137.223	125.800	130.327

Table 2 : Estimated effort of various models & proposed model.

Size(m)	Methodology	Measured Effort	Estimated Effort-Proposed Model	Bailey – Basili Estimate	Alaa F. Sheta G.E.model Estimate	Alaa F. ShetaModel 2 Estimate	Harish model1	Harish model2
2.1	28	5	4.712	7.226	8.44	11.271	6.357	4.257
3.1	26	7	6.679	8.212	11.22	14.457	8.664	7.664
4.2	19	9	8.490	9.357	14.01	19.976	11.03	13.88
12.5	27	23.9	22.811	19.16	31.098	31.686	26.252	24.702
46.5	19	79	68.190	68.243	81.257	85.007	74.602	77.452
54.5	20	90.8	77.879	80.929	91.257	94.977	84.638	86.938
67.5	29	98.4	93.437	102.175	106.707	107.254	100.329	97.679
78.6	35	98.7	106.229	120.848	119.27	118.03	113.237	107.288
90.2	30	115.8	118.753	140.82	131.898	134.011	126.334	123.134
100.8	34	138.3	130.327	159.434	143.0604	144.448	138.001	132.601



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Segmentation of Calculi from Ultrasound Kidney Images by Region Indicator with Contour Segmentation Method

By Ms.P.R.Tamilselvi, Dr.P.Thangaraj

Kongu Engineering College, Perunduri, India

Abstract - In this proposed Region Indicator with Contour Segmentation (RICS) method, five major steps are followed to select the exact calculi region from the renal calculi images. In the first and second stage, the region indices library and renal calculi region parameters are computed. After that, the image contrast is enhanced by the Histogram Equalization and the most interested pixel values of enhanced image are selected by the k-means clustering. The most interested pixel values are utilized to find the accurate calculi from the renal images. In the final stage, a number of regions are selected based on the contour process. Subsequently, pixel matching and sequence of thresholding process are performed to find the calculi. In addition, the usage of ANFIS in supervised learning has made the technique more efficient than the previous techniques. Here, the utilization of contour reduces the relative error in between the Expert radiologist and the segmented calculi, which are obtained from the proposed algorithm. Thus, the obtained error is minimized that leads to high efficiency. The implementation result shows the effectiveness of the proposed RICS segmentation method in segmenting the renal calculi in terms of sensitivity and specificity. And also, the proposed method improves the calculi area detection accuracy with reduced in computational time.

Keywords : Segmentation, Renal Calculi, Contour process, K-means clustering, ANFIS.

GJCST Classification : I.4.6



Strictly as per the compliance and regulations of:



Segmentation of Calculi from Ultrasound Kidney Images by Region Indicator with Contour Segmentation Method

Ms.P.R.Tamilselvi^a · Dr.P.Thangaraj^a

Abstract - In this proposed Region Indicator with Contour Segmentation (RICS) method, five major steps are followed to select the exact calculi region from the renal calculi images. In the first and second stage, the region indices library and renal calculi region parameters are computed. After that, the image contrast is enhanced by the Histogram Equalization and the most interested pixel values of enhanced image are selected by the k-means clustering. The most interested pixel values are utilized to find the accurate calculi from the renal images. In the final stage, a number of regions are selected based on the contour process. Subsequently, pixel matching and sequence of thresholding process are performed to find the calculi. In addition, the usage of ANFIS in supervised learning has made the technique more efficient than the previous techniques. Here, the utilization of contour reduces the relative error in between the Expert radiologist and the segmented calculi, which are obtained from the proposed algorithm. Thus, the obtained error is minimized that leads to high efficiency. The implementation result shows the effectiveness of the proposed RICS segmentation method in segmenting the renal calculi in terms of sensitivity and specificity. And also, the proposed method improves the calculi area detection accuracy with reduced in computational time.

Keywords : Segmentation, Renal Calculi, Contour process, K-means clustering, ANFIS

1. INTRODUCTION

One of the most common problems that occur in the human urinary system is renal calculi, which is often called as kidney stones or urinary stones [1]. Normally, any person affected by these kidney stone diseases will suffer from considerable pain which leads to abnormal kidney function, and also the mechanism for this disease is poorly understood so far [2]. Kidney is the most salient organ in the urinary system, which not only produce urine but also helpful in purifying the blood.

The two important functions of kidney: (i) Removing harmful substances from the blood, and (ii) Keeping the useful components in proper balance. Kidney stones appear in diverse varieties, among which the four basic types that found more often are Calcium-

containing stones, Uric acid stones, Struvite or infected stones and Cystine stones [8].

Normally the kidney diseases are classified as hereditary, congenital or acquired [14]. The detection of calcifications inside the body is a large field of study including several dynamic areas of research, which is mainly useful for diagnosing the kidney stone diseases. The actual kidney stones may be rough non-spherical in shape, but the dominant effects that are used to find the fracture in actual kidney stones, are based on the reverberation time across the length of the stone [16].

Due to the presence of powerful speckle noise and attenuated artifacts in abdominal ultrasound images, the segmentation of stones from these images is very complex and challenging [12]. Hence, this task is performed by the use of much prior information such as texture, shape, spatial location of organs and so on. Several automatic and semiautomatic methods have been proposed. Even though the performance such methods are better when the contrast-to-noise ratio is high, it deteriorates quickly when the structures are inadequately defined and have low contrast like the neuroanatomic structures, such as thalamus, globus pallidus, putamen, etc. [4]. The X-ray, positron emission tomography (PET), computer tomography (CT), Ultrasound (US) and magnetic resonance imaging (MRI) are the widely available different medical imaging modalities which are broadly employed in regular clinical practice [6]. As compared to other medical imaging modalities such as computed tomography (CT) and magnetic resonance imaging (MRI), the US is particularly difficult to segment because the quality of the image is almost low than the CT and MRI [3]. Ultrasound (US) image segmentation is greatly depends on the quality of data [7]. Moreover, it is complex to extract the features that represent the kidney tissues by segmenting the kidney region [14]. Although, ultrasound imaging is widely utilized in the medical field [13]

Ultrasound imaging is popular in the field of medicine not only due to its economical cost and non-invasive nature, but also it is a radiation-free imaging technique [12]. US imaging is economical and simple to use and also provides a faster and more exact procedures due to its real time capabilities. In numerous applications, an important role is played by the precise identification of organs or objects that are present in US

^a About : Assistant Professor, School of Computer Technology and Applications, Kongu Engineering College, Perunduri, India- 638 052
E-mail: selvipr2003@yahoo.com

^a About : Professor and Head, Department of CSE, Bannari Amman Institute of Technology, Sathyamangalam, India.
E-mail: ctpttr@yahoo.co.in

images [3]. Resolutions required by murine imaging could be achieved in ultrasonic imaging which already has a broad variety of clinical applications for human imaging, if higher frequencies (20 – 50 MHz) are used instead of the normally used frequencies (3 – 15 MHz) [5]. Speckle is a multiplicative noise, which is an important performance limiting factor in visual perception of US imaging that makes the signal or lesion complicated to identify [9, 10]. Numerous research papers have been presented on segmentation of renal calculi in US images by using diverse techniques. Since US kidney images are noisy and contain poor signal-to-noise ratio, an alternative effective techniques employing a-priori information may be utilized for compensating such problems [11]. The segmentation of renal calculi using renal images is a difficult task. Lots of researches have been performed for the successful segmentation of renal calculi using ultra sound images. A few recent related works in the literature are reviewed in the following section.

II. RELATED WORK

Benoit *et al.* [15] have proposed a region growing algorithm for segmentation of kidney stones on ureteroscopic images. Using real video images, the ground truth has been computed and the segmentation has been compared with reference segmentation. Then for comparison with ground truth, they have calculated statistics on diverse image metrics, namely Precision, Recall, and Yasnoff Measure.

Sridhar *et al.* [16] have constructed a framework for the identification of renal calculi. Normally, the kidney stones are formed by the abnormal collection of some specific chemicals such as oxalate, phosphate and uric acid. These stones can be found in the kidney, ureter or urinary bladder. Performance analysis has been performed to a set of five known algorithms by using the parameters namely success rate in calculi detection, border error metric and time. Then the best algorithm has been chosen from this performance analysis and the framework has been constructed by using this algorithm. Moreover, a procedure has been given to validate the detected calculi using the shadow that appear in ultrasound images. The algorithm has been tested by using the ultrasound images of 37 patients. The detected calculi based on the framework match those determined by professional clinicians in more than 95% of the cases.

Sridhar *et al.* [17] have developed an automated system to detect the renal calculi based on its physical characteristics. Due to the anomalous collection of certain chemicals like oxalate, phosphate and uric acid, the calculi are formed in the kidney, ureter or in urinary bladder. An algorithm has been employed to identify the calculus using its shadow. The properties of calculi such as size, shape and location have also

been extracted by their proposed system, which are crucial for reliable diagnosis. Their technique has been implemented in the MATLAB/IDL platform and a substantial success rate has been obtained.

Tamilselvi *et al.* [18] have proposed an improved seeded region growing based method which performs both segmentation and classification of kidney images with stone sizes using ultrasound kidney image for the diagnosis of stone and its early identification. The images are classified as normal, stone and early stone stage by recognizing multiple classes via intensity threshold variation diagnosis on segmented region of the images. Homogeneous region are relied on the image granularity features in the enhanced semiautomatic SRG based image segmentation process, in which the pertinent structures with dimensions similar to the speckle size extracted. The shape and size of the growing regions have relied on this look up table entries. The high frequency artifacts are also being reduced by performing region merging after the region growing. By employing the intensity threshold variation acquired for the segmented parts of the image, the diagnosis process is being performed. They have compared the size of the segmented parts of the image with the standard stone sizes i.e., if the size is below 2 mm, it is considered as absence of stone, between 2-4 mm indicates early stone stage, and 5mm & above indicates presence of kidney stones.

Tamilselvi *et al.* [19] have suggested a segmentation method for an exact segmentation of renal calculi. Classification and segmentation are the two important steps in their proposed approach. In the preprocessing stage, the image contrast improvement is being carried out by using histogram equalization and the reference pixel are selected via GA techniques before classifying a given image either as normal or stone image. The training and the classification process of diverse US images is performed by using an ANFIS system. Moreover, the same procedure is followed for the testing process of classification approach and several US images are utilized for the analysis of the precision of preprocessing classification. Subsequently, in the calculi recognition process, ANFIS is trained by using the renal calculi images having manually segmented stone regions. Several region parameters are determined and the calculi detection training process is performed by giving the result values to the ANFIS. During the testing process, the reference and testing images are compared and morphological dilation operation is applied in the calculi regions. An accurate renal calculi region was found from the result of the testing process. The experimental results have shown that their proposed segmentation method has found the accurate renal calculi from US images. They have also analyzed the performance of the proposed method by comparing it with the existing Neural Network

(NN) and SVM classifier.

The existing segmentation method has performed the calculi segmentation by region indicators and modified watershed algorithms. But in this method, the calculi detection accuracy is not satisfactory and it has produced high complexity in the calculi detection process. To avoid this drawback, we proposed a Region Indicator with Contour Segmentation (RICS) method. The outline of the paper is as follows: Section 3 briefly explains the proposed RICS segmentation process. In section 3.1, the region indicator process is explained and in section 3.2, the region parameters are computed. The contrast enhancement and most fascinated pixels by k-means clustering are explained in section 3.3 and 3.4. In section 3.5, the Contour based regions selection process is described. The experimental result and the conclusion of this paper are given in Section 4 and 5 respectively.

III. PROPOSED RENAL CALCULI SEGMENTATION TECHNIQUE

The proposed renal calculi segmentation method consists of five major steps namely, (i) Determining inner region indicators (ii) Determining the region parameters (iii) Enhancing the contrast of the image using Histogram Equalization (iv) Finding most fascinated Pixels by K-means clustering and (v) Contour based Region selection process. The proposed renal calculi segmentation training and testing procedure is shown in Figure 1.

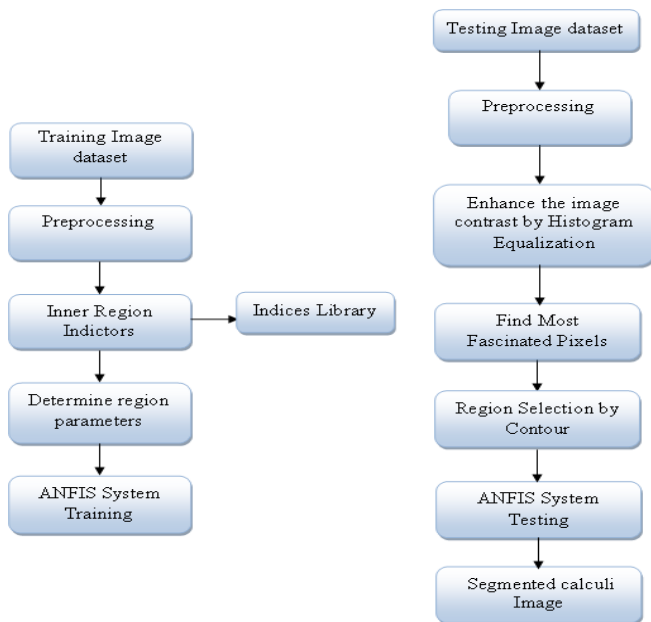


Fig. 1. Proposed RICS segmentation training and testing procedure

a) Preprocessing

In this preprocessing phase, principal component analysis with local pixel grouping (LPG-PCA)

based image denoising algorithm is used to remove the noise from the US renal calculi images.

b) Determining Inner Region Indicators

Let D represents the renal calculi image training dataset, which contains renal calculi images $D = \{I_1, I_2, \dots, I_n\}; n = 1 \dots N$, where N is the number of the renal calculi images in the given dataset D . To determine the inner region indicators, firstly, the regions representing kidney are manually marked in the known training data set ultrasound images. Then the whole image is divided into L number of blocks and for every block, an index value, which can be represented as $B = \{I_i\}; i = 1 \dots L$, is allocated. Then, each block included in B is checked to find the edge pixels present in the kidney. If any block is found to be containing edge pixels of the kidney, then the index value of the corresponding block is kept as $K = \{I_i\}; i \in L$. Hence, K which can also be called as indices library, contains the indices of blocks of all the known images.

c) Determine Region Parameters

Using the renal calculi images in D , the calculi and non calculi regions are extracted. The extracted regions from the renal calculi images are $R = \{r_1, r_2, \dots, r_m\}, m = 1 \dots M$, where M represents the total number of extracted regions. Next we find the centroids values for all the renal part of images in D , that is $C(x, y) = \{c_1^1(x, y), c_2^2(x, y), \dots, c_n^n(x, y)\}$, where $c_1^1(x, y)$ is a centroid value of image I_1 . Then, we determine the region parameters for the extracted regions from R by utilizing MATLAB function. The region parameters determined for each region are (i) Area (ii) Centroid (iii) Orientation and (iv) Bounding Box. This region parameter values are given to the ANFIS system for training process. In training process, the normal and calculi area is identified by the threshold values t_1 and t_2 . The ANFIS system result value is represented as χ . The final decision is defined by

$$\chi = \begin{cases} t_1 == \text{normal} \\ t_2 == \text{calculi} \end{cases} \quad (1)$$

d) Contrast Enhancement using Histogram Equalization

In contrast to the following enhancement process [19], initially we have converted given ultrasound image I_n^t into a grayscale image G_n^t , as histogram equalization process can be used only on

grayscale images. Histogram equalization make some enhancements to the contrast of the given gray scale ultra sound image. In histogram equalization all pixel values in gray scale image are adjusted to maximum intensity values of the image. The image that is obtained after the histogram equalization process is denoted as G_n^t .

e) *Find Most Fascinated Pixels by K-means clustering*

Mostly required pixels are computed from the image G_n^t by utilizing the k-means clustering method. K-means clustering [22] is a method of cluster analysis which aims on partition of observations into number of clusters in which each observation belongs to the cluster with the nearest mean [21]. The steps involved in the K-means clustering used in our method are described as following:-

- (i) Partition of the gray scale data points to A arbitrary centroids, one for each cluster.
- (ii) To determine new cluster centroid by calculating the mean values of all the cluster elements.
- (iii) Determining distance between the cluster centroid and the cluster elements and obtain new clusters.
- (v) Repeat process from step (i) till a defined number of iterations are performed.

The k-means algorithm aims at minimizing an objective function

$$H = \sum_{a=1}^A \sum_{g=1}^G \|d_g^a - C_a\|^2 \quad (2)$$

In eqn (2) d_g^a represents data points and C_a means center of the cluster. The resultant of the k-means clustering process has a number of clusters, which forms a cluster-enabled image I_A . Here we can select the cluster, with maximum white color pixel values, and is applied to the newly created mask I' .

f) *Contour based Region Selection Process*

Region selection process performed using renal calculi images are taken from the testing image dataset $D^t = \{I_1^t, I_2^t, \dots, I_n^t\}; n = 1 \dots N^t$, where N represents the total number of renal calculi images in the dataset D^t . The dataset D^t contains the images that are in the dimension of $P \times Q$; $1 \leq p \leq P, 1 \leq q \leq Q$. To accomplish the region selection process, a contour extraction process is utilized.

The procedure for contour based region extraction Process is as follows

Step 1

Initially the contour plot of the given gray scale image G_n^t is extracted. The contour function is described in the following equation 6. 3.

$$G_n^{tc} = (G_n^t, k) \quad (6. 3)$$

G_n^t is an input renal calculi gray scale image

- k is the number of evenly spaced contour levels in the plot
- In order to find the contour plot, the axis and their orientation and aspect ratio are defined.
- Where, G_n^{tc} represents the result of the extracted contours of renal calculi gray scale image

Step 2

After that, the final group values from the contour result image G_n^{tc} is selected. This group values contains some regions, then calculates the region parameters for that regions and the region parameters values are given to the ANFIS system that are referred in section 6.3.3.

Step 3

Then choose numbers of regions from the image G_n^{tc} which are greater than the threshold value t_1 and this selected region values are given to the empty mask S .

Step 4

The mask S contains m^s number of regions, which is represented as

$$R^s = \{r_1^s, r_2^s, \dots, r_{m^s}^s\}, m^s = 1 \dots M^s.$$

Next, compute the centroid values for the regions R^s in the mask S , it is represented as

$$C^s(x, y) = \{c_1^s(x, y), c_2^s(x, y), \dots, c_{m^s}^s(x, y)\}.$$

Step 5

There are m^s number of regions in the mask S , these mask regions are not optimal to find the exact calculi from the images. So find the optimal regions among the available regions in S by exploiting Squared Euclidean Distance (SED) between the regions.

Step 6

SED is computed between the x coordinates regions centroid values $C^s(x, y)$ and training images

centroid values $C(x, y)$ and y coordinates regions centroid values $C^S(x, y)$ and training images centroid values $C(x, y)$ values individually.

Step 7

The SED difference process is described in the following equ.2&3 for both x and y coordinates values.

$$\phi(x) = (c_1^I(x) - c_1^S(x))^2 + (c_2^I(x) - c_2^S(x))^2 + \dots + (c_n^I(x) - c_{m_s}^S(x))^2 \quad (2)$$

$$\phi(y) = (c_1^I(y) - c_1^S(y))^2 + (c_2^I(y) - c_2^S(y))^2 + \dots + (c_n^I(y) - c_{m_s}^S(y))^2 \quad (3)$$

The values $\phi(x)$ and $\phi(y)$ are compared with the threshold value t_3 . If any one of the result values $\phi(x)$ or $\phi(y)$ are greater than the given threshold value t_3 , that corresponding region are selected.

Step 8

Then, the last group values are selected from the contour method result image G_n^{tc} and have placed these values into the newly created mask M . After getting the result from contour process, the pixel matching and Multidirectional traversal operation is performed.

Pixels Matching: Here, first step is to divide the mask image M into m number of blocks and the index values $I^x = \{I_m^x\}$ are allocated. Then $I^x = \{I_m^x\}$ is compared against K by using the following conditions

- (i) Retain the pixel values in the block $m \in M$; if an index value $I_m^x = I_L$, then.
- (ii) Change the block $m \in M$ pixel values into 0, or else

And hence M' is generated. Over M' and I' an AND operation is performed followed by a morphological dilation operation and hence the resultant image U is obtained.

Multidirectional Traversal: Here we have proposed two major traversals called bottom-up traversal and top-down traversal. In each of the traversal, a left-right traversal is applied. The traversals are applied over U , which is binary. At the time of two major traversals, once the pixel with '1' is obtained, then left-right traversal is enabled so that all the regions in the same axis and the region of the first obtained pixel are removed from the mask. The survived pixel values are marked into the original test image and it is subjected to the consequent process of Thresholding.

Thresholding: Here, a chain of thresholding process is performed in the original image.

- Firstly, the pixel values that are marked by using the previous process are compared against a

defined threshold value t_3 . The pixel values those are greater than t_3 are stored in a newly created mask U_s .

- The region parameters are determined for the regions in the mask U_s and the computed region parameters are given to the ANFIS to obtain the ANFIS score. If the ANFIS score is greater as compared to t_4 , then the selection of regions is performed.
- Then, we count the number of neighbor pixel values around the selected regions which are greater than the threshold value t_5 , and the number of count value of each region is compared with the threshold value t_6 . If the count value is greater than the threshold value t_6 , then the regions are selected.
- The selected regions from the previous thresholding process are involved in the morphological dilation operation. After the morphological operation, count the number of regions that are presented in the image. The region count value is compared with two threshold values t_7 and t_8 .
- If the count value is greater than t_7 , then perform the traversing down operation once, and if the count value is greater than t_8 , then perform traversing down operation in multiple times.
- In the final thresholding process, each regions area value is calculated and it is compared with the threshold values t_9 and t_{10} . The regions that are less than t_9 and greater than t_{10} are selected. The selected regions are then placed into the original testing image I_n^t .

By performing all the above described process in various renal calculi kidney images, the calculi region is segmented.

IV. RESULTS AND DISCUSSION

The proposed RIC segmentation technique is implemented in MATLAB platform (version 7.10) and the performance of the proposed RIC segmentation method is evaluated using 50 images. In the proposed RICS segmentation method, five major steps are performed over these training and testing renal calculi and renal ultra sound images. The sample input normal and calculi images are shown in figure 2.

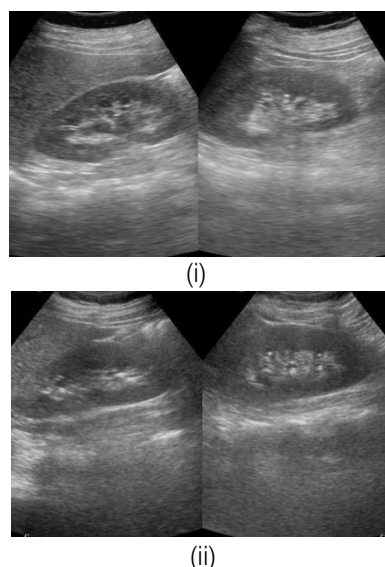


Figure 2: Sample Input Renal Images (i) Normal Renal Image (ii) Renal Image with Calculi

The region parameter values are computed for the 110 training images and these parameters result values are given to the ANFIS system to perform the training process. The region parameter values are well trained in the ANFIS system and this performance is evaluated with testing renal calculi images. 50 testing images are involved in the testing process. Figure 3 shows the result of the histogram equalization, k-means clustering and contour method.

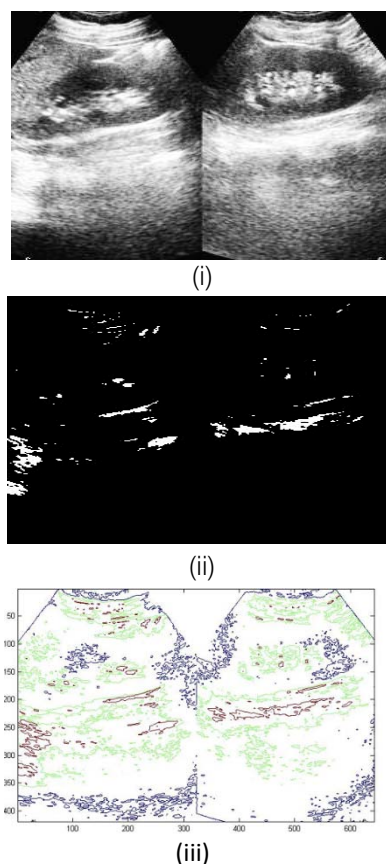


Figure 3: Result images are obtained from (i) Histogram Equalization (ii) K-Means clustering (iii) Contour Method

The histogram equalized image contrast is enhanced when compared to the original input image shown in Fig 2 (i). In Fig. 3 (ii), the same pixel values are grouped into number of clusters values and this can be used to find the most interested pixel values. The result images in fig 3 (iii) shows that the contour method has divided the testing image pixels into three groups by representing three different colors. The selected group value from the contour method result is shown in figure 4.

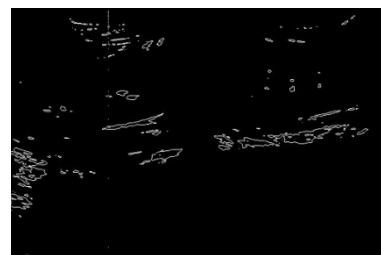
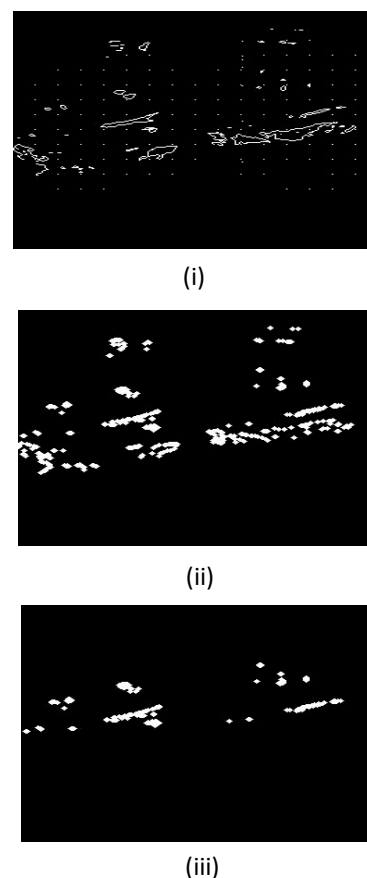
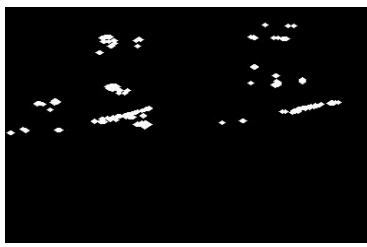


Figure 4 : Selected Group Value Result from the Contour Method

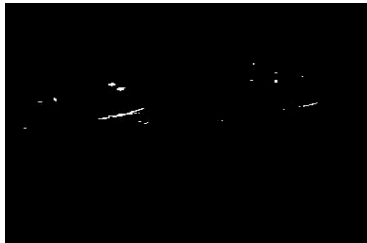
After the contour process, the chain of dilation and traversing operation are performed in the processed renal calculi image results that are shown in figure 5. The traversing operations eliminate the most unwanted regions from the renal calculi images, so as to easily find the calculi from the image. Subsequently, the thresholding process intermediate results are shown in the figure 6.





(iv)

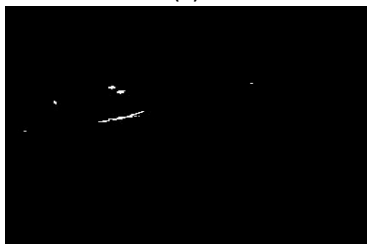
Figure 5: Result Images from (i) Pixel Matching (ii) Morphological Dilation (iii) Traversing up and (iv) Traversing down



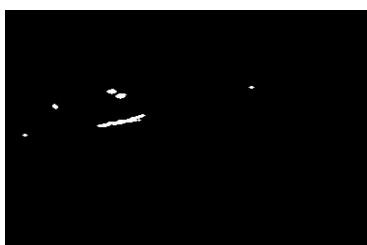
(i)



(ii)



(iii)



(iv)



(v)

Figure 6 : Thresholding Process Results

Finally, the selected regions from the thresholding process are given to the original image that is demonstrated in the following figure 7. In figure 7, the calculi regions are exactly marked in red color. The result image has shown that the proposed RIC segmentation method has exactly found the calculi region from the renal calculi images. The performance of proposed RIC segmentation method is analyzed with different images and it is described in the following section.

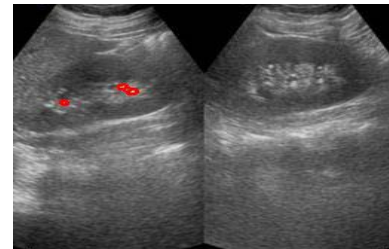


Figure 7 : Proposed RIC Segmentation Result Image

a) Performance Analysis

The performance of the RICS segmentation method by using four testing images is given in table 1. This performance analysis exploits statistical measures [20], to compute the accuracy of calculi segmentation done by the RIC segmentation method. The performance of the RIC segmentation analysis is shown in the below Table 1.

ID No	Se	Sp	Acc	FPR
1	93.33	99.93	99.92	0.07
2	89.06	100	99.98	0
3	100	100	100	0
4	100	100	100	0
Average	95.60	99.98	99.98	0.02

ID No	PPV	NPV	FDR	MCC
1	59.96	99.99	40.04	74.77
2	100	99.98	0	94.36
3	100	100	0	100
4	100	100	0	100
Average	89.99	99.99	10.01	92.28

Table 1: Statistical performance measures for four different US renal calculi

In Table 1, we have achieved high sensitivity, specificity and accuracy level in 1 sec computational time. The segmented stone area by RICS segmentation method is compared with previous IORM segmentation method and conventional segmentation algorithms. A

relative error is calculated between the segmented stone area marked by the expert radiologist and the proposed method. The formula for the calculation of relative error is described below,

$$\nu = \left| \frac{E - P}{E} \right| \times 100 \quad (4)$$

Where, ν - Relative Error

E - Stone area marked by Expert radiologist

P - Stone area marked by the proposed RICS segmentation method

The stone area marked by the expert radiologist, the RICS segmentation method and its relative error are given in Table 2.

Expert radiologist (mm ²)	Stone area (mm ²)	Relative error of RICS method
88.4	88.3	0.113
61.5	61.6	0.163
64.0	64.0	0.000
91.9	90.3	1.741
72.4	72.4	0.000
25.0	24.9	0.400
121.7	121.1	0.493
36.0	36.0	0.000

Table 2 : RICS segmentation relative error performance

The computational time of RICS method is obtained from the calculi detection process. The average Computational time of the system is shown in the Table 3 for four renal calculi images. The computational time of RICS method is very low .

Systems	Images	Computational Time(sec)
RICS	1	1.455874
	2	1.572855
	3	1.807253
	4	1.685209
	Average	1.630298

Table 3 : Computational time of the RICS Methods

V. CONCLUSION

In this paper, a RICS segmentation method to segment the calculi from the renal calculi images was

proposed. The proposed method was implemented and set of renal calculi images were utilized to evaluate the proposed RICS segmentation method. The proposed method has exactly detected the calculi and produced a high segmentation accuracy result. The performance of RICS segmentation method was analyzed has produced less relative error. Moreover, our proposed RICS segmentation method has produced 99.98% of accuracy, 95.60% sensitivity and 99.98 % specificity values.

REFERENCES REFERENCES REFERENCIAS

- Ioannis Manousakas, Chih-Ching Lai and Wan-Yi Chang, "A 3D Ultrasound Renal Calculi Fragmentation Image Analysis System for Extracorporeal Shock Wave Lithotripsy", International Symposium on Computer, Communication, Control and Automation, Vol. 1, pp. 303-306, 2010
- Jie-Yu He, Sui-Ping Deng and Jian-Ming Ouyang, "Morphology, Particle Size Distribution, Aggregation, and Crystal Phase of Nanocrystallites in the Urine of Healthy Persons and Lithogenic Patients, " IEEE Transactions On Nanobioscience, Vol. 9, No. 2, pp. 156-163, June 2010
- Jun Xie, Yifeng Jiang and Hung-tat Tsui, "Segmentation of Kidney From Ultrasound Images Based on Texture and Shape Priors", IEEE Transactions On Medical Imaging, Vol. 24, No. 1, pp. 45-56, January 2005
- Ujjwal Maulik, "Medical Image Segmentation Using Genetic Algorithms", IEEE Transactions On Information Technology In Biomedicine, Vol. 13, No. 2, pp. 166-173, March 2009
- Erwan Jouannot, Jean-Paul Duong-Van-Huyen, Khalil Bourahla, Pascal Laugier, Martine Lelievre-Pegorier and Lori Bridal, "Comparison and Validation of High Frequency Ultrasound Detection Techniques in a Mouse Model for Renal Tumors", In Proceedings of IEEE International Ultrasonics Symposium, Vol. 1, pp. 748-751, 2004
- Mahmoud Ramze Rezaee, Pieter M. J. van der Zwet, Boudewijn P. F. Lelieveldt, Rob J. van der Geest and Johan H. C. Reiber, "A Multiresolution Image Segmentation Technique Based on Pyramidal Segmentation and Fuzzy Clustering", IEEE Transactions On Image Processing, Vol. 9, No. 7, pp. 1238-1248, July 2000
- Alison Noble and Djamal Boukerroui, "Ultrasound Image Segmentation: A Survey", IEEE Transactions on Medical Imaging, Vol. 25, No. 8, pp. 987-1010, August 2006
- Saurin R. Shah, Manhar D. Desai and Lalit Panchal, "Identification of Content Descriptive Parameters for Classification of Renal Calculi", International Journal of Signal and Image Processing, Vol.1, No. 4, pp.

- 255-259, 2010
9. C.P.Loizou, C.Christodoulou, C.S.Pattischis, R.S.H.Istepanian, M.Pantziaris and A.Nicolaides, "Speckle Reduction in Ultrasound Images of Atherosclerotic Carotid Plaque", In Proceedings of IEEE International Conference on Digital Signal Processing, Santorini, Greece, Vol. 2, pp. 525-528, 2002
10. C.P.Loizou, C.S.Pattischis, R.S.H.Istepanian, M.Pantziaris, T.Tyllis and A.Nicolaides, "Quality Evaluation of Ultrasound Imaging in the Carotid Artery", In Proceedings of IEEE Mediterranean Electro technical Conference, Dubrovnik-Croatia, Vol. 1, pp. 395 – 398, 2004
11. Bommanna Raja, Madheswaran and Thyagarajah, "A General Segmentation Scheme for Contouring Kidney Region in Ultrasound Kidney Images using Improved Higher Order Spline Interpolation", International Journal of Biological and Life Sciences, Vol. 2, No. 2, pp. 81-88, 2006
12. Abhinav Gupta, Bhuvan Gosain and Sunanda Kaushal, "A Comparison of Two Algorithms for Automated Stone Detection in Clinical B-Mode Ultrasound Images Of the Abdomen", Journal of Clinical Monitoring and Computing, Vol. 24, pp. 341–362, 2010
13. Ratha Jeyalakshmi and Ramar Kadarkarai, "Segmentation and feature extraction of fluid-filled uterine fibroid-A knowledge-based approach ", Maejo International Journal of Science and Technology, Vol. 4, No. 3, pp. 405-416, 2010
14. K. Bommanna Raja · M. Madheswaran and K. Thyagarajah, "Texture pattern analysis of kidney tissues for disorder identification and classification using dominant Gabor wavelet", Journal Machine Vision and Applications, Vol. 21, No. 3, pp. 287-300, 2010
15. Neil R. Owen, Oleg A. Sapozhnikov, Michael R. Bailey, Leonid Trusov and Lawrence A. Crum, "Use of acoustic scattering to monitor kidney stone fragmentation during shock wave lithotripsy", In Proceedings of IEEE International Ultrasonics Symposium, Vancouver, Canada, pp. 736-739, 2006
16. Ioannis Manousakas, Yong-Ren Pu, Chien-Chen Chang and Shen-Min Liang, "Ultrasound Image Analysis for Renal Stone Tracking During Extracorporeal Shock Wave Lithotripsy", In Proceedings of IEEE EMBS Annual International Conference, New York City, USA, pp. 2746-2749, 2006
17. Neil R. Owen, Michael R. Bailey, Lawrence A. Crum and Oleg A. Sapozhnikov, "Identification of Kidney Stone Fragmentation in Shock Wave Lithotripsy", In Proceedings of IEEE Ultrasonics Symposium, New York, NY, pp. 323-326, 2007
18. Tamilselvi and Thangaraj, "Computer Aided Diagnosis System for Stone Detection and Early Detection of Kidney Stones", Journal of Computer Science, Vol. 7, No. 2, pp. 250-254, 2011
19. P. R. Tamilselvi and P. Thangaraj, "An Efficient Segmentation of Calculi from US Renal Calculi Images Using ANFIS System", European Journal of Scientific Research, Vol. 55, No.2, pp. 323-333, 2011
20. [20]http://en.wikipedia.org/wiki/Sensitivity_and_specificity
21. Kehar Singh, Dimple Malik and Naveen Sharma, "Evolving Limitations in K-Means Algorithm in Data Mining and their Removal", International Journal of Computational Engineering & Management, Vol. 12, pp. 105-109, 2011





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

An Efficient and Speedy Activity Model for Information System Based Organizations

By Md. Nazrul Islam, Mahmuda Akter, Dr. M.A. Kashem, Md. Mostafizer Rahman

Dhaka University of Engineering and Technology

Abstract - This paper presents an activity model that addresses the responsibilities among different C-level Officers in any IT-reliant systems in organizations. The activity model provides an integrated set of actions that extend and clarify the work system framework and related work system concepts, thereby helping in understanding, analyzing, and designing technical and sociotechnical systems. The activity model is an advance step toward an enhanced work system approach that is quickly accessible, understandable and clear visualization to business professionals, is more rigorous than most current applications of work system concepts, and can be linked more directly to precise, highly detailed analysis and design approaches for IT professionals. Specification of the activity model clarifies ambiguities in the work system framework and forms a clearer conceptual basis for tools and methods that could improve communication and collaboration between business and IT professionals through e-Media.

Keywords : Activity model, Information Systems, Work System Framework, C-level Officers, Business Professionals, Work System Environment.

GJCST Classification : D.2.9



Strictly as per the compliance and regulations of:



An Efficient and Speedy Activity Model for Information System Based Organizations

Md. Nazrul Islam^α, Mahmuda Akter^Ω, Dr. M.A. Kashem^β, Md. Mostafizer Rahman^ψ

Abstract - This paper presents an activity model that addresses the responsibilities among different C-level Officers in any IT-reliant systems in organizations. The activity model provides an integrated set of actions that extend and clarify the work system framework and related work system concepts, thereby helping in understanding, analyzing, and designing technical and sociotechnical systems. The activity model is an advance step toward an enhanced work system approach that is quickly accessible, understandable and clear visualization to business professionals, is more rigorous than most current applications of work system concepts, and can be linked more directly to precise, highly detailed analysis and design approaches for IT professionals. Specification of the activity model clarifies ambiguities in the work system framework and forms a clearer conceptual basis for tools and methods that could improve communication and collaboration between business and IT professionals through e-Media.

Keywords : Activity model, Information Systems, Work System Framework, C-level Officers, Business Professionals, Work System Environment.

I. INTRODUCTION

Unfinished and run-away projects, Workforce ignoring, Information overload, Employee mistrust, Security breaches all are the challenges of any IS based Work System due to poorly aligned with business executives. User requirements, the costs required to implement a new technology or update a technology dose not always realized by the C-level Executives because of communication and technology knowledge gap in organization. This paper addresses two aspects of the problem: 1) intangible and procedural gaps among different C-levels and 2) communication and knowledge gaps that blocks and degrade collaboration between business and IT professionals which also effects planning and

Implementation in any business organization. E-Media which plays an important role in business organization can develop the relationship among C-level. The aim of this paper is to describe an activity model for describing, analyzing, and designing both corporate business system and IS together in a Work System. (Activity model are graphical representations of workflows of stepwise activities and actions with support for choice, iteration and concurrency.[1] In the Unified Modeling Language, activity model can be used to describe the business and operational step-by-step workflows of components in a system. An activity shows the overall flow of control that is conditional or parallel. For example, the activity flow from customer to Chief Marketing Officer after an action Request for a new/updated technology is performed by the customer then CMO could perform another action (Verify requests) in any Work system.) The activity model is a significant reformulation and extension of the work system framework, which has a number of limitations that the model addresses.

The rest of this paper is prepared as follows segment II shortly illustrates the Background, Segment III described literature review, Segment IV described Model requirements, Segment V includes work system framework, Segment VI includes proposed work system environment, Segment VII establishes an activity model, Segment VIII describes the elements of an activity model, Segment IX includes implementation phase and lastly Segment X concludes the paper.

II. BACKGROUND

This segment provides background about the work system approach in general and the work system framework in particular. A work system approach assumes that the unit of analysis is a work system, a business system in which human participants and/or machines perform work (processes and activities) using information, technology, and other resources to produce specific products and/or services for specific internal or external customers. Work systems change over time through iterations of planned change (projects) and through incremental adaptations and innovations that may be unplanned. New technology, government regulations, agri-terrorism, and biological threats are forcing work system to change the way they face these challenges. As a result, work system and going towards with information technology in the same alignment. So a

*Author^α : Assistant Professor in Computer Science and Engineering (CSE) in Dhaka University of Engineering and Technology (DUET), Gazipur-1700, Bangladesh. Ph: +88 01732183690
E-mail : nazrul_ruet@yahoo.com*

*Author^Ω : B.Sc. in engineering degree in Information and Communication Technology (ICT) from Mawlana Bhashani Science and Technology University, Bangladesh. Ph.: +88 01723720525
E-mail : badhan_mbstu@yahoo.com*

*Author^β : Associate Professor in Computer Science and Engineering (CSE) in Dhaka University of Engineering and Technology (DUET), Gazipur-1700, Bangladesh. Ph: +88 01199117213
E-mail : drkashemll@yahoo.com*

*Author^ψ : Assistant Programmer in Dhaka University of Engineering and Technology (DUET), Gazipur-1700, Bangladesh. He has interest in Wireless Networking. Ph: +88 01719472181
E-mail : mostafiz26@gmail.com*

better work flow model e.g., activity model, where collaboration is maintained between the business professionals and information executives to achieve goals within time frame.

III. LITERATURE REVIEW

Identifying a model is not a new topic in business organization or work system or academic research. However, activity flow on technology based work system has not been heavily researched in either the academic or public forums which leads to success rate on IS/IT projects remains unacceptably low. Much of the previous literature has focused on only their operations but poor business/IT communication, poor user participation in projects, lack of support by business executives, difficulties with implementation in organizations, technical and conceptual complexity of IS/IT projects, poor resources for projects, unrealistic project schedules, and staff turnover. IS/IT research has addressed these issues from various viewpoints, such as studying:

- A. business/IT communication and business/IT alignment (e.g., Reynolds and Yetton (2009)).
- B. the usefulness and pitfalls of IS development tools for IT professionals (e.g., zur Muhlen and Recker, 2008;)
- C. concepts and models related to paths to success (e.g., Value Chain Analysis, Critical Success Factor models, Business System Planning (BSP), Strategic System Planning (SSP), the technology acceptance model)
- D. reducing gap between sociotechnical and technical views of systems (e.g., metamodel, Steven Alter, 2010)

Finally, this paper's activity model fits into basic research concerning IS/IT concepts. If reducing communication and knowledge gaps were an important goal, the education of IT professionals might recognize more fully that most business professionals are more concerned with improving business performance rather than with specifying IT based tools (such as UML, Microsoft Visio) that they might use and get accurate design methodology.

IV. MODEL REQUIREMENTS

Using the Laudon & Laudon definition of information systems the core requirements of an information system exists to collect, process, store and distribute information that supports the decision making and control of an organization. This provided the first requirement of our model. Our model must be able to identify the mechanics necessary to gather, process, store, and distribute information. Second the model must also identify if the information gathered is processed to support management of the organizations. Establishing these criteria as a means for identifying the existence information systems requires that our model

be two pronged. The first prong must be able to identify both the information processes and the reason for gathering information and the second prong must identify that the information decision making, coordination, and control in an organization. In addition to supporting being processed supports management before existence of an information system can be verified.

V. WORK SYSTEM FRAMEWORK

Alter (2002) defines, a "work system is a system in which human participants and/or machines perform business processes using information, technologies, and other resources to produce products and/or services for internal or external customers". Alter's work system framework is developed from nine core elements as displayed in Fig.1. The first four elements: information, participants, processes, and technologies, constitute the systems doing the work. These first four elements define what Alter refers to as a basic system within his framework. The work systems output are the products and services received by its customers. The remaining three elements: environment, strategies, and infrastructure influence the overall process to determine "if a work system can operate as intended and can accomplish its goals" (Alter, 2002). Alter's framework provides a two-step approach to explore the activities established in an organization. The first step, the work system, identifies the components of the systems doing the work. The second step identifies the interaction of the work system with the environment, strategies, and infrastructure. This second step provides data on how systems are used which brings us to the second prong of our model development.

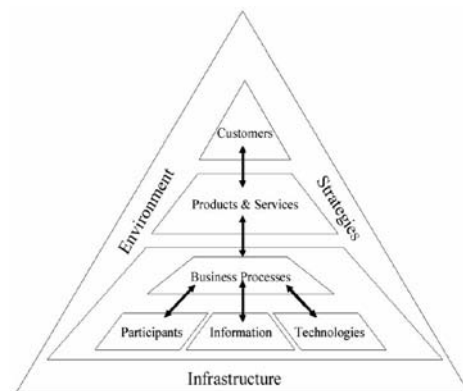


Fig.1 : Alter's Work System

Possible alternative frameworks. The work system framework was developed over time to guide its users to develop a basic understanding of an IT-reliant work system in an organization. However alternative frameworks are:

- A. Business process. Consists of only one element instead of the nine elements of the work system framework. The work system approach has been called a business process approach, but it involves

- much more than just the detailed logic of the business process.
- B. Input-processing-output. Effective for computer operation but less useful for describing the operation of IT-reliant work systems in organizations.
 - C. People, process, technology. 3-sided framework in three boxes is a reminder that people, process, and technology are relevant for thinking about systems in organizations.
 - D. SIPOC. A 5-element model used in Six Sigma analysis: supplier, input, processing, output, customer but does not clearly specified.
 - E. GRITIKA ontology. An ontology containing 7 concepts: goal, role, interaction, task, information, knowledge, and agent. Suggested by Zhang et al. (2004) for modeling e-service applications.
 - F. Zachman (2008) framework. A 6X6 framework outlining an enterprise's architecture, and therefore at a different level than a work system model. The 6 rows include scope, business model, system model, technology model, detailed representations, and functioning enterprise. The six columns include what, how, where, who, when, why.
 - G. Steven Alter's (2010) metamodel. an integrated set of concepts can be described using its 31 elements that bridge the chasm between sociotechnical and technical views of systems in organizations. But approach is mechanistic, does not focus on process of change, work flows, triggering conditions, resource requirements, business rules, and post-conditions of specific activities.

- C. e-Banking Online banking (or Internet banking) allows customers to conduct financial transactions on a secure website operated by their retail or virtual bank, credit union or building society.
- D. Other Work System is mentioned here a competitor who analysis provides both an offensive and defensive strategic context to identify opportunities and threats.
- E. Economic Indicators (or business indicator) here is a statistic about the economy who allows analysis of economic performance and predictions of future performance of a work system.
- F. e-Learning covers a wide set of applications and processes, such as Web-based learning, computer-based learning, virtual meeting, and digital collaboration. It includes the delivery of content via Internet, intranet/extranet (LAN/WAN), audio- and videotape, satellite broadcast, interactive TV, and CD-ROM.

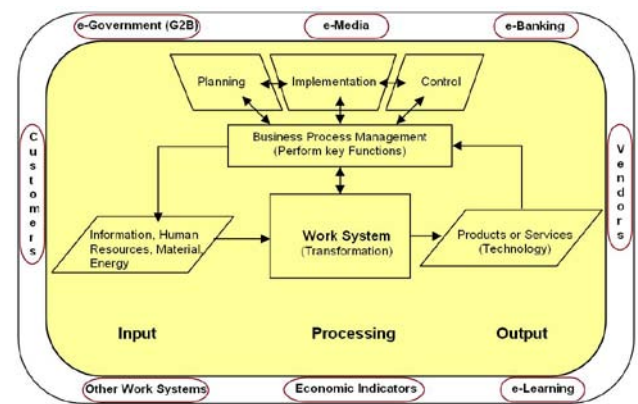


Fig.2 : The work System Environment

Alter argues for an IS as a special type of work system. A work system is a system in which humans and/or machines perform work using resources to produce specific products and/or services for customers. An IS is a work system whose activities are devoted to processing (capturing, transmitting, storing, retrieving, manipulating and displaying) information [5]. The environment of any information based work system encompasses of input, processing and output. The work system environment is the "suprasystem" within which an organization operates and often determines how a system must function. As shown in Figure-2, the work system environment consisting of e-Governance (G2B), e-Media, e-Bank, Customers, Vendors, Other Competitive Work System, Economic Indicators, e-Learning will provide constraints and consequently, influence the actual performance of the work system can be described as follows:

- A. e-Government (short for electronic government) here is a digital interactions between a government businesses/Commerce (G2B)
- B. e_Media (short of Electronic media) are in the form of digital media known as video recordings, audio recordings, multimedia presentations, slide presentations, CD-ROM and online content.

Kay's (Kay & Edwards, 1999) functions of Work System can be expressed as a cycle where information is used to navigate and move through each of the functions: planning, implementation, and control. This cycle is also illustrated in Fig.2.

Planning is the first function also referred to as the strategic decision stage to identify problems or strategic direction occurs. The Business Process Management must identify the problem or opportunity and choose to act or not. The second function is implementation which is selecting and acting on a plan. Once a plan is identified and approved, resources and infrastructures are put into place with the technical side. Progress is verified on a routine basis to determine if the actions put into place are moving the Business Process Management towards the intended goals. Control or monitoring the progress of an action plan is the last function. If the progress is not acceptable, that is the invented technology does not follow the standard of a disqualified product then the overall process in work may be terminated by the Business Process Management.

VII. ESTABLISHING AN ACTIVITY MODEL FOR WORK SYSTEM

Fig.3. is an activity model for the analysis, design and implementation of new or updated technology in information system based organization. Actual work flow from worker to different C-level executives and related necessary actions done by the business process management is easily understandable from the figure. As earlier mentioned most of the work system consists of Kay's (Kay & Edwards, 1999) functions like planning, implementation and control. So we divide the work system into three phases according to Kay's function. Then activity flow among the phases is easily understandable by subdividing the phases.

In planning phase, Customer plays an important role as because most requests e.g., new or updated technology in an organization comes from customer. All sorts of requests are gathered and noted by the Chief Marketing Officer (CMO) who has investigated primarily on the market demand. Then the requests are sending to the corporate officers (CEO, COO) to take necessary actions to fulfill the customer demand. A meeting will be called to collaborate with all C-level officers to gather

different ideas, methodologies, related necessary technologies, budget to build a plan. The plan may be discarded by the Corporate Officers if it will fail the budget or fall in time limit exceed. By the end of this phase the budget is passed for the request to implement otherwise.

In the implementation phase, the Corporate Officer will asked the Technical Officer (CTO, CIO) to implement IS/IT strategies. Emerging all resources e.g. IT-Participants, Non-IT-Participants, Software, Technological Entity, Hardware with all Infrastructure e.g., Human Infrastructure, IT Researcher/IT Specialists, Project Manager, Programmer, Technical and Information develop a new or updated technology to fulfill the user requests. In this stage a primary verification is done on the new born or updated technology. Circular verification is done by changing parameters like resources, infrastructure, methodologies, ideas, tools and techniques if it does not fulfill user requirements. In this phase IT-Participants can share their acceptance or demand with the existing technologies to the Technical Officers. Implementation phase will turn off if the technology verification is okay.

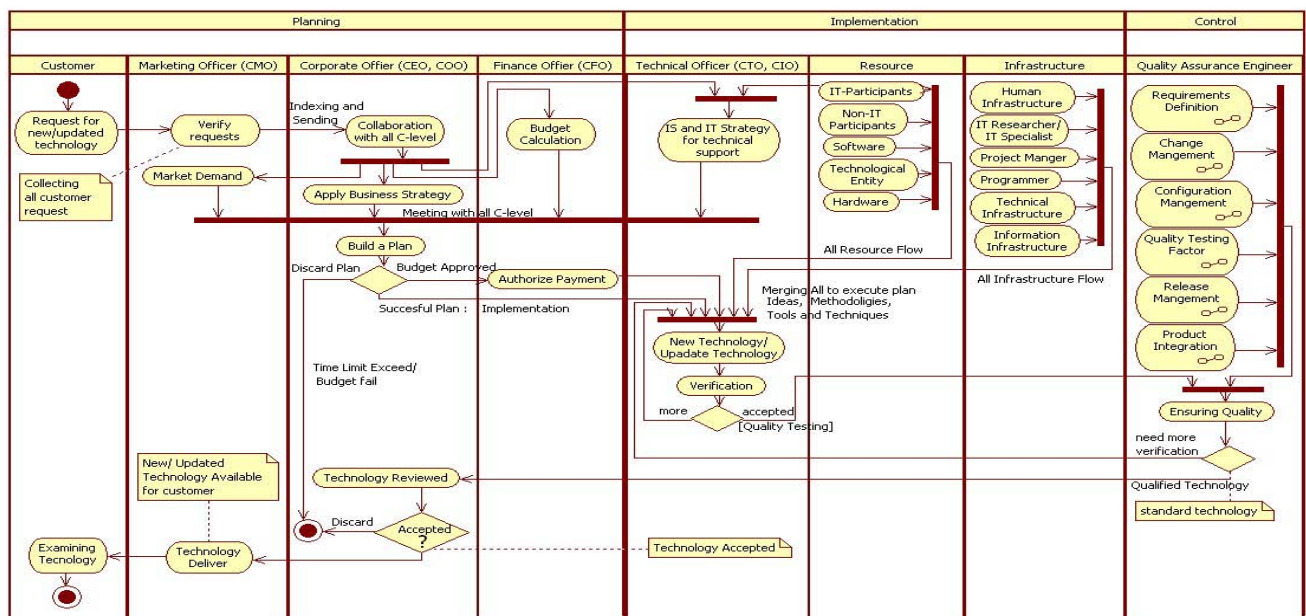


Fig.3 : Activity Model for IS based Work System

In the final phase of Kay's function control, a series of sub-activity like Requirements Definition, Change Management, Configuration Management, Quality Testing Factors, Release Management and Product Integration is performed together on the technology or products to test its quality with the standards. If it does not meet the quality parameters well, some sort of more verification is needed in this

stage. A note for the qualified product is maintained for future learning in this phase. Then the product is returned to the business process management. Then again planning if fails the corporate level satisfaction or stop planning is corporate officers accept and order the CMO to deliver the new/updated technology or product. Thus a qualified product or technology goes to the customer for examining their requirements level.

VIII. ELEMENTS OF THE ACTIVITY MODEL

- A. **Work systems functioning:** The proposed work systems can be functioned by three phases like planning, implementation and control to execute its goal.
- B. **Infrastructure.** Includes relevant human, information, and technical resources that are used by the work system but are managed outside of it and are shared with other work systems.
- C. **Customers.** Are requester of the work system' technology and recipients of products and services for purposes other than performing work activities within the work system.
- D. **Marketing Officer.** To whom demand of new or updated technology request arises according to market demand.
- E. **Corporate Officer.** Who makes decision and execute a plan, permit to implement and finally controlled or monitor the overall work system.
- F. **Strategies.** Are relevant to a work system include enterprise strategy, organization strategy, and work system strategy. In general, they are the business policies to achieve the work system's goal. Here business strategies and IS/IT strategies differ from the point of technology related.
- G. **Finance Officer.** Permit budget according to higher authority.
- H. **Technical Officer.** A brief knowledge about all technical, operational and informational gather together to implement and verify permitted plan according to very fast upcoming technology.
- I. **Resource.** Are needed to implement work system's plan including Participants, Software, Technological Entity and Hardware. Participants are people who perform working within the work system, including both IT-Participants and Non-IT Participants. IT-Participants can share their expectations from technology to the technical officers.
- J. **Infrastructure.** Include Human, IT Researcher/IT Specialists, Project Manger, Programmer, Technical and Information. Here Human plays a minor role like clerical work in the work system and Information is what we used, created, captured, transmitted, stored, retrieved, manipulated, updated, displayed, and/or deleted by a specific activity in the activity model.
- K. **Quality Assurance Engineer.** Measure the quality of a product by composing some sub-activity like Requirements Definition, Change Management, Configuration Management, Quality Testing Factors, Release Management and Product Integration. Always perform a comparison against standard.
- L. **Products and services:** consist of information, physical things, and/or actions produced by a work system in a word here it is called Technology.

IX. MODEL IMPLEMENTATION

Through the use of the model illustrated in Figure-3 it was determined that each work system implemented the mechanics of their system in a unique fashion to achieve their individual business process management goals. The complexity of any work system can easily be optimized by importing and applying the activity model in any information system based organization. The activity model proved successful in determining the activity flow in the work system in passing some action state and three phases like planning, implementation and control. By completing a cycle of the activity model a new or refined technology or product will outcome.

X. CONCLUSION

The activity model of Unified Modeling Language is used in teaching and research, helping corporate executives to learn information technology, to maintain collaboration among C-levels, to give value of technological ideas and work, to distribute work flow among different levels, to monitor overall work. The activity model spells out the shortcomings of the all work system framework. The research is going on, so activity model is a standard for any work system that always integrates information system to their business process to reveal, enlarge and enriches the work system with technologies.

REFERENCES REFERENCES REFERENCIAS

1. Alter, S. The Work System Method: Connecting People, Processes, and IT for Business Results. Works System Press, CA.
2. Alter, S. (2002). The work system method for understanding information systems and information system research. Communications of the Association for Information Systems, 9, 90-104.
3. Alter, Steven, "BRIDGING THE CHASM BETWEEN SOCIOTECHNICAL AND TECHNICAL VIEWS OF SYSTEMS IN ORGANIZATIONS" (2010). ICIS 2010 Proceedings. Paper 54.
4. Böhm, Markus; Nominacher, Bastian; Fähring, Jens; Leimeister, Jan Marco; Yetton, Philip; and Krcmar, Helmut, "IT CHALLENGES IN M&A TRANSACTIONS – THE IT CARVE-OUT VIEW ON DIVESTMENTS" (2010). ICIS 2010 Proceedings. Paper 105.
5. Jean S. Adams, 2009. "Establishing a Model to Identify Information Systems in Nontraditional Organizations". Information Systems Education Journal, 7 (88).
6. Kay, R.D. and Edwards, W. M. (1999). Farm Management, 4th ed. McGraw-Hill, Boston.
7. Keen, P. G. W. (1978). Decision support systems: an organizational perspective. Reading, Mass., Addison-Wesley Pub. Co. ISBN 0-201-03667-3.



8. Kim, Heekyung H. and Brynjolfsson, Erik, "CEO Compensation and Information Technology" (2009). ICIS 2009 Proceedings. Paper 38.
9. Leavitt, H.J. 1965. "Applied organisational change in industry: structural, technological, and humanistic approaches," pp. 1144-1170 in March, J.G. (ed.) Handbook of organizations. Chicago: Rand McNally.
10. Zhang, H., Kishore, R. Ramesh, R. and Sharman, R. 2004. "The GRITIKA Ontology for Modeling e-Service Applications: Formal Specification and Illustration," Proceedings of the 37th Hawaii International Conference on System Sciences.
11. Zachman, J.A. 2008. "The Zachman FrameworkTM: The Official Concise Definition," Zachman International enterprise architecture, <http://www.zachmaninternational.com/index.php/home-article/13#maincol>.





GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Computer Vision Based Collaborative Augmented Reality Method for Human-Computer Interaction

By Akhil Khare, Vinaya Kulkarni, Dr. Akhilesh Upadhayay

Abstract - Computer vision is becoming very popular now a days since it can hold a lot of information at a very low cost. With this increasing popularity of computer vision there is a rapid development in the field of virtual reality as it provides an easy and efficient virtual interface between human and computer. At the same time much research is going on to provide more natural interface for human-computer interaction with the power of computer vision .the most powerful and natural interface for human-computer interaction is the hand gesture. Hand replaces the currently used cumbersome and inefficient devices like mouse and keyboard and with the bare hands one can easily communicate with the computer. This paper explores a system where hand gesture can be effectively used as a password in the login process for authentication of a person using just a simple web camera. Also this technique does not need any special device like head-mounted display, gloves or any special camera that operates beyond visible spectrum. So with this idea, with a simple video camera and bare hands, a person can interact with computer.

Keywords : Computer vision, Human-Computer interaction, Gesture recognition, Haar-like feature.

GJCST Classification : I.2.10



A COMPUTER VISION BASED COLLABORATIVE AUGMENTED REALITY METHOD FOR HUMAN-COMPUTER INTERACTION

Strictly as per the compliance and regulations of:



A Computer Vision Based Collaborative Augmented Reality Method for Human-Computer Interaction

Akhil Khare^α, Vinaya Kulkarni^Ω, Dr. Akhilesh Upadhayay^β

Abstract - Computer vision is becoming very popular now a days since it can hold a lot of information at a very low cost. With this increasing popularity of computer vision there is a rapid development in the field of virtual reality as it provides an easy and efficient virtual interface between human and computer. At the same time much research is going on to provide more natural interface for human-computer interaction with the power of computer vision .the most powerful and natural interface for human-computer interaction is the hand gesture. Hand replaces the currently used cumbersome and inefficient devices like mouse and keyboard and with the bare hands one can easily communicate with the computer. This paper explores a system where hand gesture can be effectively used as a password in the login process for authentication of a person using just a simple web camera. Also this technique does not need any special device like head-mounted display, gloves or any special camera that operates beyond visible spectrum. So with this idea, with a simple video camera and bare hands, a person can interact with computer.

Keywords : Computer vision, Human-Computer interaction, Gesture recognition, Haar-like feature

1. INTRODUCTION

With the rapid increase in human and computer interaction an easy and natural interface is getting much more value than it was previously. Now a day's keyboard and mouse are used as the main interface for transferring information and commands to the computer. In our day to day life we human uses our vision and hearing as a main source of information about our environment. Therefore a much research is going on for providing more natural interface for human-computer interaction based on computer vision. Hand gesture is most popular and effective medium for communication in virtual environment because it conveys information very effectively and naturally. The purpose of this project is to develop new perceptual interfaces for human-computer-interaction based on visual input captured by computer vision systems. In the initial days different cumbersome devices were imposed on users such as head mounted display, digital gloves etc. These devices had limited the users movement and feels uncomfortable to the user. On the other hand vision based gesture recognition system that uses bare hand is becoming very popular because it does not need any device to impose on user's body Instead, it provides a natural hand gesture recognition interface system for human-computer interaction. The whole

process of hand gesture recognition is broadly divided in three steps first is the segmentation that is the hand is separated from the background using different methods such as colour segmentation method. Then the features of the hand are extracted that is the feature detection and with the help of extracted features multiple hand gestures are categorized in to three groups communicative, manipulative and controlling gesture. Communicative gesture is used to express an idea or concept. Manipulative gesture is used to interact with virtual objects in virtual environment. To control a system controlling gestures are used.

Previous methods suffers from the limitation of lightening changes and less accuracy. Also in some methods different devices were used such as head mounted display or hand gloves etc. In some methods two cameras were used as a well as sometimes a 3D sensor was also used. So a new method has been invented for gesture recognition that uses haar like structure along with topological features and color segmentation technique to identify and classify different hand postures and gives satisfactory performance with higher accuracy when applied to human-computer interaction for personal authentication. This method makes use of a single camera to capture the image as well as no special device or sensor is needed.

In this paper we focus our attention to vision based recognition of hand gesture for personal authentication where hand gesture is used as a password. Different hand gestures are used as password for different personals. This method could also be used for blind people who can use their hand gesture as a password for the login process. Hand gesture has been used mostly to convey some commands to the computer. This system is introduces a new application of hand gesture that is the personal authentication.

The remainder of this paper is structured as follows: section II takes a short review on different methods described in various papers. The hand gesture classification and phases are discussed in section III. Section IV discusses proposed system and also discusses how it is different from the existing systems and finally the conclusion.

II. LITERATURE ANALYSIS

One of the method proposed by Rokade et al [1] uses the technique of thinning of segmented image, but it needs more computation time to detect different hand postures. One method is based on elastic graph matching, but it is sensitive to light changes [2]. In a system proposed by Stergiopoulou and Papamarkos YCbCr color segmentation model was used but the background should be plane and uniform [3].

In one method CSS features were used by Chin-Chen Chang for hand posture recognition [4]. In the method presented by this paper a boost cascade of classifiers trained by Adaboost and haar like features are used to accelerate the evaluation speed used to recognize two hand postures for human-robot recognition. It uses haar like features along with color segmentation method to improve the accuracy in detecting the hand region and then the topological method is used to classify different hand postures.

In the method proposed by Shuying Zhao [5] for hand segmentation Gaussian distribution model (for building complexion model) is used. With Fourier descriptor and BP neural network an improved algorithm is developed that has good describing ability and good self learning ability. This method is flexible and realistic. In the system proposed by Wei Du and Hua Li statistic based and contour based features are used for stable hand detection [6].

In a system developed by Utsumi [7] a simple hand model is constructed from reliable image features. This system uses four cameras for gesture recognition. In a system known as fingermouse developed by Quek the hand gesture replaces mouse for certain actions [8]. In this system only one gesture that is pointing gesture is defined and for mouse press button shift key on the keyboard is used. Segan has developed a system [9] that uses two cameras to recognize three gestures and hand tracking in 3D. By extracting the feature points on hand contour the thumb finger and pointing finger are detected by the system.

In the system presented by Triesch multiple dcues such as motion cue stereo cue, color cue are used for robust gesture recognition algorithm[10]. This system is used in human robot interaction that helps robot to grasp objects kept on the table. In the system real time multihand posture recognition system for human-robot interaction haar like feature and topological features were used along with color segmentation technique [11]. This method gives accurate results and a rich set of features could be extracted.

Compared with the traditional interaction approaches, such as keyboard, mouse, pen, etc, vision based hand interaction is more natural and efficient. Yikai Fang, Kongqiao Wang, Jian Cheng and Hanqing Lu proposed a robust real time hand gesture recognition

method [12] in their paper "A real time hand gesture recognition method". In this method, firstly, a specific gesture is required to trigger the hand detection followed by tracking; then hand is segmented using motion and color cues; finally, in order to break the limitation of aspect ratio encountered in most of learning based hand gesture method, the scale-space feature detection is integrated into gesture recognition. Applying the proposed method to navigation of image browsing, experimental results show that this method achieves satisfactory performance.

Wei Du and Hua Li presented a real-time system in "Vision based gesture recognition system with single camera" for human-computer interaction through gesture recognition and hand tracking[10]. Stable detection can be achieved by extracting two kinds of features: statistic-based feature and contour-based feature. Unlike most of previous works, our system recognizes hand gesture with just one camera, thus avoids the problem of matching image features between different views. This system can serve as a natural and convenient user input device, replacing mouse and trackball.

III. GESTURE CLASSIFICATION AND MODELLING

Hand gesture is a movement of hand(s) and arm(s) that are used as a means to express an idea or to convey a command to control an action. Hand gesture can be classified in a several ways. For HCI applications the most commonly used and suitable classification divides hand gesture in to three groups communicative, Manipulative and controlling gestures. To express an idea or a specific concept communicative gestures are used. It may be used as a substitute for verbal communication. Communicative gesture is similar to sign language and as in sign language it also requires a high structured set of gestures. To interact with objects in an environment manipulative gestures are used. This is generally used to manipulate virtual objects in virtual environment. Controlling gestures as the name indicates used to control a system or to locate an object. One of the application of controlling gestures is controlling mouse movements on the desktop. The major steps in hand gesture analysis are analysis of hand motion, modeling of hand(s) and arm(s), mapping the motion features to the model and interpreting the gesture in the time interval. But generally speaking analysis of hand gesture is totally application dependent.

From the psychological point of view hand gesture consists of three phases, these are preparation,, nucleus and retraction phase. Preparatory phase includes bringing the hand from resting position to the starting posture of the gesture. Sometimes this phase is very short and many times it is combined with the

retraction phase of the previous gesture. The next is nucleus phase that includes the main concept of gesture and has a definite form and that is used as a command to the computer. The last phase is retraction that shows the resting position of hand after completing the gesture. If the gesture is succeeded by another gesture then the retraction phase may be very short or not present. The preparatory and retraction phases are usually and hand movements are faster compared to the nucleus phase. However identifying starting and ending position of the nucleus phase is quite difficult as there are variations in the preparatory and retraction phase.

IV. PROPOSED SYSTEM

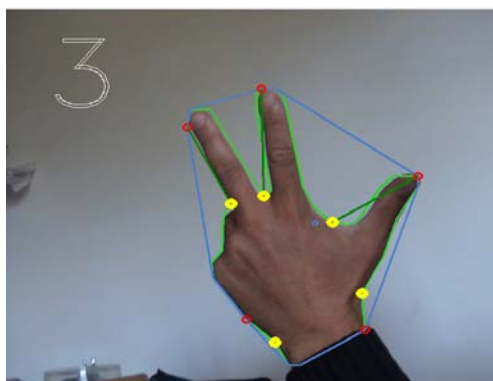


Fig 1(a)



Fig 1(b)

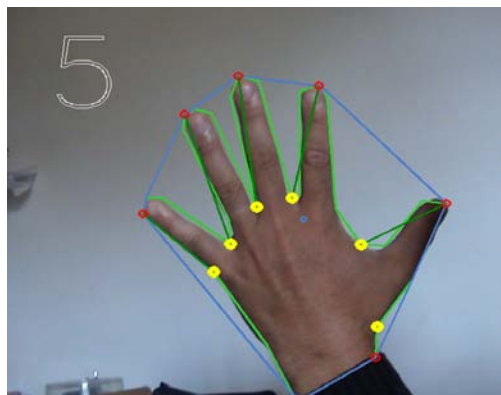


Fig 2(a)



Fig 2(b)

Fig 1 : (a) and 2(a) original picture and fig 1(b) and 2(b) segmentation result

a) Video Capturing

Video capture is the process of converting an analog video signal—such as that produced by a video camera—to digital form. The resulting digital data are referred to as a digital video stream, or more often, simply video stream. This is in contrast with screen casting, in which previously digitized video is captured while displayed on a digital monitor.

The video capture process involves several processing steps. First the analog video signal is digitized by an analog-to-digital converter to produce a raw, digital data stream. In the case of composite video, the luminance and chrominance are then separated; this is not necessary for S-Video sources. Next, the chrominance is demodulated to produce color difference video data. At this point, the data may be modified so as to adjust brightness, contrast, saturation and hue. Finally, the data is transformed by a color space converter to generate data in conformance with any of several color space standards, such as RGB and YCbCr. Together, these steps constituted video decoding, because they "decode" an analog video format such as NTSC or PAL. Special electronic circuitry is required to capture video from analog video sources. At the system level this function is typically performed by a dedicated video capture card. Such cards often utilize video decoder integrated circuits to implement the video decoding process.

b) Image extraction from video

Here we have to select captured video as input. We are now ready to start extracting frames from the videos. After getting frame from video start to extract images from those frames. Store those extracted files in particular folder.

c) Image enhancement and Remove noise

Noise reduction is the process of removing noise from a signal. Noise reduction techniques are conceptually very similar regardless of the signal being processed, however a priori knowledge of the characteristics of an expected signal can mean the

implementations of these techniques vary greatly depending on the type of signal. The median filter is a nonlinear digital filtering technique, often used to remove noise. Such noise reduction is a typical pre-processing step to improve the results of later processing (for example, edge detection on an image). Median filtering is very widely used in digital image processing because under certain conditions, it preserves edges while removing noise. The main idea of the median filter is to run through the signal entry by entry, replacing each entry with the median of neighboring entries. The pattern of neighbors is called the "window", which slides, entry by entry, over the entire signal. For 1D signal, the most obvious window is just the first few preceding and following entries, whereas for 2D (or higher-dimensional) signals such as images, more complex window patterns are possible (such as "box" or "cross" patterns). Note that if the window has an odd number of entries, then the median is simple to define: it is just the middle value after all the entries in the window are sorted numerically.

d) Background suppress

An algorithm that detects and removes background shadows from images in which the pattern set occupies the upper-most intensity range of the image and the image is background dominant outside the pattern set is presented. The algorithm will remove background shadows and preserve any remaining texture left behind by the shadow function. A mathematical model of the histogram modification function of the shadow-removal algorithm is developed. An analysis of the sequential nature of the algorithm is included along with simulated results to verify the mathematical model developed and to show the effectiveness of the algorithm in removing background pattern shadows.

e) Hand region segmentation

The initial step of hand gesture recognition is the detection of hand region from the background. This step is also known as hand detection. It involves detecting and extracting hand region from background and segmentation of hand image. Previous methods made use of following two approaches that is the color based model and statistical based model. This system uses the additional third approach i.e. haar like feature with adaboost technology. Different features such as skin colour, shape, motion and anatomical models of hand are used in different methods. The output of this step is a binary image in which skin pixels have value 1 and non-skin pixels have value 0.

Haar-like detector: First step is conversion of input image to an integral image since haar-like features can be calculated from an integral image with a greater speed. A rich set of haar like features can be computed from the integral image. The integral image at the point $p(x,y)$ contains the sum of the pixel values left and above this point. It is defined as,

$$P(x,y) = \sum I(X,Y)$$

$$x' \leq x, y' \leq y$$

each haar like feature is composed of two connected black and white rectangles. The value of a haar likefeature is obtained by subtracting the sum pixel values of the white rectangle from the black rectangle. Single haar like feature is not able to recognize hand region with high accuracy. The adaboost learning algorithm can considerably improve the overall accuracy stage by stage by using a linear combination of these individually weak classifiers. This combination makes the processing faster and robust.

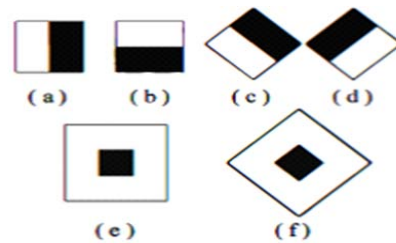


Fig.3 : Haar-like features

Different colour models can be used for hand detection such as YCbCr, RGB, YUV etc. but the proposed system uses YCbCr color model since it is useful in compression applications. Also YCbCr separates RGB in to luminance & chrominance information. Following equation is used to transform RGB values in to YCbCr color space'. The characteristics of hand shape such as topological features could be used for hand detection. Learning detectors from pixel values: Hands can be found from their appearance and structure such as Adaboost algorithm. 3D model based detection: Using multiple 3D hand models multiple hand postures can be estimated.

f) Feature extraction and gesture recognition

The next important step is hand tracking and feature extraction. Tracking means finding frame to frame correspondence of the segmented hand image to understand the hand movement. Following are some of the techniques for hand tracking.

- Template based tracking: If images are acquired frequently enough hand can be tracked. It uses correlation based template matching. by comparing and correlating hand in different pictures it could be tracked.
- Optimal estimation technique: Hands are tracked from multiple cameras to obtain a 3D hand image.
- Tracking based on mean shift algorithm: To characterize the object of interest it uses color distribution and spatial gradient. Mean shift algorithm is used to track skin color area of human hand.

Two types of features are there first one is global statistical features such as centre of gravity and second one is contour based feature that is local feature that includes fingertips and finger-roots. Both of these features are used to increase the robustness of the system. Hand posture can be distinguished using the number of fingers of the hand and if the number of fingers are same then the angle between two fingers can be measured to recognize the specific gesture.

The goal of hand gesture recognition is interpretation of the meaning gesture of the hands location and posture conveys. From the extracted features multiple hand gestures are recognized. Different methods for hand gesture recognition can be used such as template matching, method based on principle component analysis, Boosting contour and silhouette matching, model based recognition methods, Hidden Markov Model (HMM). Hand gesture is movement of hands and arms used to express an idea or to convey some message or to instruct for an action. From psychological point of view hand gesture has three phases.

g) Register user

The Register User action registers the user information with the installer to identify the user of a product. it provides a unique user id for every user. a large set of postures and gestures is stored on the computer one for each individual.

h) Login

When a user wants to login he/she has to perform the desired hand gesture. This hand gesture can be performed using single hand. That gesture will be compared with the already recorded gesture that works as a password for that particular person, if that gesture matches with the performed gesture then only that person will be authenticated and will be allowed to access his/her account or product. Basic idea is that the number of fingers are counted and the password is created Ex, 123,432,531,23,4532,123451 etc. the password can be any combination of the numbers from 0,1,2,3,4,5. This password performed by the user is authenticated by the system and he/she will be allowed to access the application or is rejected the access.

This proposed system could be used by any application to authenticate the authorized user. The major benefit of this system is that it could be used by blind users also, but the accuracy is the major concern, the system may not give accurate results in intricate background.

V. CONCLUSION

Vision based hand gesture recognition has major applications in human-computer interaction as well as in intelligent service robot. This paper describes a collaborative vision based hand gesture recognition system where a hand gesture could be effectively used by a person as her password for the personal

authentication. This system provides an easy interface for human-computer interaction. This system will provide a more efficient system with greater accuracy that makes use of both the hands as well as the drawback of previous techniques have been tried to remove such as complexion problem could be effectively removed by using background model alongwith the complexion model. In multihand gesture recognition method a rich set of features could be extracted using haar like feature and topological features with greater accuracy.

REFERENCES REFERENCES REFERENCIAS

1. R. Rokade, D. Doye, and M. Kokare, "Hand Gesture Recognition by Thinning Method," Digital Image Processing, 2009 IEEE International Conference on, 2009, pp. 284-287.
2. Stergiopoulou, N. Papamarkos, "Hand gesture recognition using a neural network shape fitting technique," Engineering Applications of Artificial Intelligence, Vol. 22, Issue 8, December 2009, pp. 1141-1158.
3. C.C. Chang, "Adaptive multiple sets of css features for hand posture recognition," Neurocomputing, vol. 69,2006, pp.2017-2025.
4. Shuying Zhao, Wenjun Tan, Chengdong Wu, Chunjiang Liu, Shiguang Wen, A Novel Interactive Method of Virtual Reality System Based on Hand Gesture recognition" 2009 Chinese Control and Decision Conference (CCDC 2009).
5. Wei Du Hua Li, "Vision based gesture recognition system with single camera", proceedings of ICSP2000, 2000 IEEE.
6. Akira Utsumi, Tsutomu Miyasato, Fumio Kishino, Ryohei Nakatsu, Hand Gesture Recognition System using Multiple Cameras, Proceedings of ICPR'96, vol 1~~6 6 7 - 6 7 A1u, wst 1996.
7. Francis K.H. Quek, Non-Verbal Vision-Based Interfaces, Vision Interfaces and Systems Laboratory Technical Report 1994, Electrical Engineering and Computer Science Department, The University of Illinois at Chicago.
8. Jakub Segen and Senthil Kumar, Human-Computer Interaction using gesture Recognition and 3D Hand Tracking, Proceedings of ICIP'98, Vol. 3, pp188-192, Chicago, October 1998.
9. Jochen Triesch, Christoph von der Malsburg, Robotic Gesture Recognition, Proceedings of ECCV'98, pp 2 3 3 -2 2 4 1998.
10. Cao Chuqing, Li Ruifeng, Ge Lianzheng"Real time multihand posture recognition" , International Conference On Computer Design And Appliations (ICCD A 2010)
11. Yikai Fang, Kongqiao Wang, Jian Cheng, Hanqing Lu, "A Real-Time Hand Gesture Recognition Method," Multimedia and Expo, 2007 IEEE International Conference on, July 2007, pp.995-998.



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Headlight Prefetching for Cooperative Media Streaming in Mobile Environments

By M. Divya, Vasanth Sena, T.P. Shekar, N. Sathish Kumar

Sri Chaitanya College of Engineering, karimnagar

Abstract - Multimedia information services in mobile environments are becoming more and more important with the proliferation of technologies. Media streaming, in particular, is a promising technology for providing services such as news clips, live sports. To avoid service interruption when the users keep moving, proper data management strategies must be employed. We propose a new headlight prefetching technique for the streaming access points to deal with the uncertainty of client movement and the requirement of seamless service handoff. For each mobile client, we maintain a virtual fan shaped prefetching zone along the direction of movement similar to the headlight of a moving vehicle. The overlapping area and the accumulated virtual illuminance of the headlight zone on a particular cell determines the degree and volume of prefetching to be made by the streaming access point of that cell. Headlight prefetching solves the issues of identifying the streaming access points responsible for prefetching, the timing and the amount of data to prefetch in a single mechanism which is simple and effective. Simulation results demonstrate that our techniques can significantly decrease streaming disruptions, reduce bandwidth consumption, increase cache utilization and improve service response time.

GJCST Classification : C.2.m



HEADLIGHT PREFETCHING FOR COOPERATIVE MEDIA STREAMING IN MOBILE ENVIRONMENTS

Strictly as per the compliance and regulations of:



Headlight Prefetching for Cooperative Media Streaming in Mobile Environments

M. Divya^a, Vasanth Sena^a, T.P. Shekar^b, N. Sathish Kumar^ψ

Abstract - Multimedia information services in mobile environments are becoming more and more important with the proliferation of technologies. Media streaming, in particular, is a promising technology for providing services such as news clips, live sports. To avoid service interruption when the users keep moving, proper data management strategies must be employed. We propose a new headlight prefetching technique for the streaming access points to deal with the uncertainty of client movement and the requirement of seamless service handoff. For each mobile client, we maintain a virtual fan shaped prefetching zone along the direction of movement similar to the headlight of a moving vehicle. The overlapping area and the accumulated virtual illuminance of the headlight zone on a particular cell determines the degree and volume of prefetching to be made by the streaming access point of that cell. Headlight prefetching solves the issues of identifying the streaming access points responsible for prefetching, the timing and the amount of data to prefetch in a single mechanism which is simple and effective. Simulation results demonstrate that our techniques can significantly decrease streaming disruptions, reduce bandwidth consumption, increase cache utilization and improve service response time. To avoid disconnection and/or service breakdown when the users keep moving, proper data management strategies must be taken by all parties. We propose a twolevel framework and a set of new techniques for cooperative media streaming in mobile environments.

I. INTRODUCTION

Media streaming is a promising technology for providing multimedia information services such as news clips, live sports, and hot movies in mobile environments. Effective data management for media streaming is naturally the key to the successfulness of such services. Since many users may request the same media, traditional client server model can easily result in server bottleneck, bandwidth waste, poor cache utilization and longer delay. Furthermore, the mobility of mobile users raise the issue of seamless service hand off. To avoid service interruption, proper data management strategies must be employed. In this paper, we propose a new technique named headlight prefetching for media streaming in mobile environments. The technique is designed for the streaming access

points to deal with the uncertainty of client movement, the unpredictability of request pattern and the requirement of seamless service handoff. For each mobile client, we maintain a virtual fan shaped prefetching zone along the direction of movement similar to the headlight of a moving vehicle. The overlapping area and the accumulated virtual illuminance of the headlight zone on a particular cell determines the degree and volume of prefetching to be made by the streaming access point of that cell. When a mobile user makes an unexpected sharp turn, the headlight shifting technique is used such that all the previously prefetched media segments can be easily shifted to accommodate the new direction of movement. For users requesting the same media at around the same time, the headlight sharing technique is developed for sharing the headlight zones to avoid repeated prefetching. The set of techniques solve the issues of identifying the streaming access points responsible for prefetching, the timing and the amount of data to prefetch in a single mechanism which is simple, intuitively appealing and effective.

To evaluate the effectiveness of our techniques, we construct a Java based simulation environment and compare the performance of different combinations of our schemes with on demand techniques. Simulation results demonstrate that, for streaming media services in mobile environments, headlight prefetching, shifting and sharing are simple and effective techniques which can significantly decrease streaming disruptions, reduce bandwidth consumption, increase cache utilization and improve service response time.

The rest of the paper provides a survey of related issues and research work and presents the media streaming system infrastructure we assume and characterize the challenges of dynamic data management in such environments. We also introduces the idea of headlight prefetching and the associated data management techniques. we outline the simulation environment for experimenting with our ideas and present the results of performance evaluation on different aspects of the prefetching techniques.

II. STREAMING MEDIA SERVICE ARCHITECTURE

For media streaming services in mobile environments, we assume a system architecture. The multimedia information are provided by streaming media

*Author^a : Sri Chaitanya College of Engineering, karimnagar(AP).
E-mail : Divya_162001@gmail.com*

Author^a : Asso. Professor, Sri Chaitanya College of Engineering, karimnagar(AP). E-mail : vasanthaizza@yahoo.com

Author^b : Asso. Professor, Sri Chaitanya College of Engineering, karimnagar(AP). E-mail : tpshekhhar@gmail.com

Author^ψ : Asso. Professor, SVS Group of Institutions, Hanamkonda, Warangal(AP). E-mail : nskumar.svs@gmail.com

servers(SMSs). All cells have corresponding streaming access points(SAPs) connected with each other through traditional fixed link networks. Each SAP provides wireless media services for mobile users within that cell. Users not reachable from any SAP are disconnected. In general, let the local wireless access cost between a mobile user and an SAP be L , and the remote access cost of requesting the unit from an SMS be R . Then, based on the nature of the architecture and service charge, we assume that L is larger than R since wireless access cost is usually larger than that of fixed link. However, if a media segment is available locally on an SAP, then it only takes a local access cost of L to service the segment. If the segment is not available on the SAP cache, then a remote access cost of R must be added in addition to the local access cost. Therefore, a remote access is always more expensive than a local access. Since intermittent disconnection is unavoidable in mobile environments, effective data management strategies must be employed to conserve access cost and reduce play interruption.

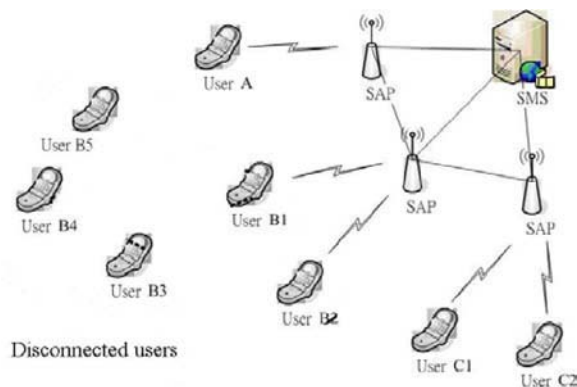


Fig 3.1 : System architecture for mobile streaming service

A streaming media is usually considered as a sequence of media segments. Because of the streaming nature, it is not necessary to provide the entire media all at once. A media request from a mobile client can therefore be treated as the starting request of a sequence of media segments. There are several data management challenges for streaming media services in mobile environments:

- A mobile user can request any media from any location at any time. On a request, the SAP of the cell where the user resides must locate and send the first segment as soon as possible to reduce service delay.
- Once started, the SAP must fetch and transmit the subsequent segments fast enough to catch up with the playing speed. Good prefetching and buffering techniques are required to avoid possible interruption.
- A mobile user can move and change direction at any time. Such a dynamism can only be handled by close coordination of neighboring SAPs to provide

seamless streaming media services across cell boundaries.

- To reduce cost, the media segments should be served on a proximity basis. In other words, it is best to use the local SAP or to locate a nearby SAP with the requested segments to provide the service. The remote access to the SMS should only be used as a last resort.

Our goal is to develop effective dynamic data management techniques to answer the challenges of streaming media services in mobile environments. In particular, the headlight prefetching and associated techniques are designed to solve the problems of identifying responsible SAPs, determining prefetching segment assignment and handling the dynamic data management issues in a simple and unified framework. Our schemes are unique in several respects:

- The idea of virtual fan shaped headlight prefetching zone is simple, intuitive appealing and efficient. We can use the headlight coverage area to identify the prefetching SAPs and the virtual illuminance to determine the look ahead window for each SAP. Both are straightforward and easy to compute.
- Headlight prefetching is flexible and dynamic adjustable. We can use the shape and size of the virtual fan to control the range and degree of prefetching. For mostly straight and fast moving users (eg. vehicles on a freeway, passengers on a train, ...), we can have a smaller central angle and longer radius. For slow moving clients that tend to wonder around, a larger central angle and smaller radius allow more neighboring SAPs to be prepared for serving the clients.
- Should a mobile user make an unexpected sharp turn, we provide the headlight shifting technique such that all the previously prefetched segments can be easily shifted to the SAPs along the new direction.
- To efficiently reuse prefetched segments, the headlight sharing technique is developed to coordinate neighboring SAPs on media services and to avoid repeated segment prefetching.

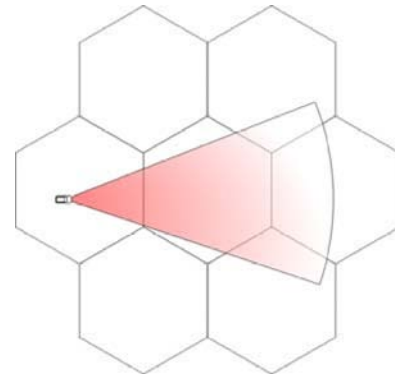
streaming media services in mobile environments entail significant challenges on highly dynamic and efficient data management. The mobile and streaming characteristics of the services also provide a window of opportunity for information system designers. We take on this opportunity to provide simple and effective dynamic data management techniques that adhere closely to the movement patterns of mobile users. It turns out that high quality streaming media services in mobile environments can indeed be achieved with proper coordinations of SAPs and right strategies for data management.

III. HEADLIGHT PREFETCHING

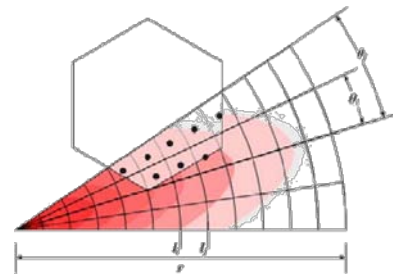
As stated earlier, an important characteristic of streaming media is the continuous playing requirement. Once a media is started, the requesting user expects a smooth and seamless viewing experience. Fetching a segment upon its request is not likely to catch up with the playing speed. Prefetching and buffering are almost a must in such case. Traditional prefetching is simply done by maintaining a sliding window immediately ahead of the current segment by the SAP in charge of the media service. This is only satisfactory when the user is moving in a straight ahead pattern. For irregular moving patterns, the traditional approach may fail miserably. The problem is that if the user changes direction and moves toward a new SAP, the later may not have anything prepared for the user. The uncertainty of user movement can easily result in frequent disruption and unpleasant viewing experience. To have all surrounding SAPs prefetch the needed segments for the user is clearly not cost effective. We therefore need a good mechanism to continuously identify the proper set of prefetching SAPs for a moving client. Those SAPs that are more likely to be visited next should be selected with higher probability. Even if the set of prefetching SAPs is successfully identified, we still have another important problem of determining the right media segments for each SAP to prefetch. The simplest approach of having all selected SAPs prefetch the same set of media segments is certainly not satisfactory. To save processing and communication costs, the ideal case is to prefetch just the segments needed to keep the media viewing smooth. Similar to the previous issue, we need to cope with the uncertainty of user's moving speed and pattern by dynamically determining the prefetching segments assignment. Those SAPs that are more likely to be visited next should be assigned more segments to prefetch. In addition to the prefetching SAPs identification and segment assignment problems, we also need to determine the right timing for the SAPs to start or stop prefetching. Otherwise the prefetched segments may either arrive too early or too late.

The idea of headlight prefetching is to have a simple and unified mechanism for solving the issues discussed above. A headlight prefetching zone is a virtual fan shaped area along the direction of user movement similar to the headlight of a moving vehicle. All SAPs of the cells touched by the headlight zone are selected as the prefetching SAPs for the user. The overlapping area and the accumulated virtual illuminance of the head light zone on a particular cell determines the degree and volume of prefetching to be made by the SAP of the cell.

Headlight prefetching solves the issues of prefetching SAPs identification, segment assignment and the prefetch timing in a single mechanism which is simple and intuitive appealing.



The headlight prefetching zone.



The headlight model.

The headlight prefetching zone is modeled by two parameters. The radius r determines the extent of look ahead for prefetching. The angle θ is used to control the span of coverage. Both are dynamically adjustable. In general, faster moving users need longer radiuses. Users that tend to wonder around need larger θ s to have more SAPs ready to carry on the services when a user changes direction unexpectedly. The headlight zone serves as a prediction of possible future interaction of the user with neighboring cells. By using such a simple and intuitive appealing metaphor, we can easily identify the set of SAPs that need to prefetch media segments for the user.

Once the prefetching SAPs identification problem is solved, we still need to take on the segment assignment problem. The head light metaphor gives us yet another simple way to handle this problem. A basic characteristic of a vehicle headlight is that the farther away from the vehicle the lower the illuminance. The area immediately in front of the vehicle has the highest brightness. This characteristic matches exactly the requirement of the segment assignment problem. Since a user can change direction at any time, the media segments prefetched by the SAPs farther away from the user are less likely to be actually used. They should be assigned fewer segments to save the cost. On the other hand, the SAPs closer to the user are more likely to be responsible for providing the media services. They

should be allocated more segments to prevent undesirable disruption. Due to such a close resemblance, we use the accumulated virtual illuminance of the cell area covered by the headlight zone as the weight for segment assignment. In this way, we solve both problems in a unified framework.

The headlight zone and the computation of virtual illuminance.

The problem is in determining the cell area covered by the headlight zone since the intersected area could be in any shape. To avoid the costly computation of the covered area, we use a much simpler approach to approximate the virtual illuminance. More specifically, we partition the headlight zone into smaller grids and precompute the virtual illuminance as well as the center of each grid. Since each grid is in a regular shape, the virtual illuminance can be easily computed by the following formula.

On the need to determine the segment assignment of a particular cell, we simply add up the illuminance of all grids whose centers fall inside the cell. Since the granularity of the grids can be changed easily, we can have higher level of precision at any time by using finer partitioning.

a) Prefetching Segment Assignment

To determine the segment assignment, we need to consider both the user movement and media playing speed. Therefore the number of segments that should be handled by the current cell is tP . If the current media segment being played is S_i , then $SAPC$ must prefetch the segments $S_{i+1}, S_{i+2}, \dots, S_{i+tP}$. The first segment that need to be prefetched by $SAPC_3$ is S_{i+tP+1} . The expected number of segments to be handled is t_3P . The starting segment and the total number of segments to prefetch for all other $SAPs$ within the headlight zone can be determined in a similar way. Since it is clearly not cost worthy for all such $SAPs$ to prefetch the full range of segments, therefore we use the virtual illuminance as a weighting factor to determine the actual number of segments to prefetch. More specifically, the exact segment assignment for $SAPC''$ to prefetch is Parameters for determining the segment assignment

b) Headlight shifting

Headlight prefetching is quite effective for largely stable moving users. However, if the user makes a sharp turn, then most or even all of the prefetching done on the previous headlight zone may be completely useless since the user is no longer heading toward the predicted direction. We can of course start a new round of head light prefetching for the new situation. However, this will double the prefetching cost. Since the media stream is played continuously, the segments needed for the new headlight zone overlap significantly with that of the old zone. Therefore the better way is to shift these segments to the new zone. We call the idea headlight

shifting. For efficient headlight shifting, we need to solve at least two problems:

- Since the same media segment may be available on more than one $SAPs$, how to choose or map the shifting source and target?
- Different $SAPs$ may have different number of prefetched segments available for shifting, how do we distribute and balance the shifting load?

We propose three strategies for headlight shifting. The simplest and most intuitive way is to map the grids of the new zone to corresponding grids of the old zone with possible offset based on user's current position. This is called direct mapping. More specifically, let the old and new grids be $G_{i,j}$ and $G_{3i,j}$ respectively. Also assume that the position where the user makes the turn is in grid $G_{x,y}$. Then we simply map $G_{3i,j}$ to $G_{i+x,j}$. In other words, the SAP whose cell covers $G_{3i,j}$ can simply request the media segments to be directly shifted from the SAP whose cell covers $G_{i+x,j}$.

Direct mapping works fine if the user is at a position close to the entrance of a grid since most of the segments are not yet played and therefore the mapping is useful and efficient. However, the problem with direct mapping is that if the user is at a position about to leave a grid then the mapping is likely to be off by one grid since most of the segments have already been played. The problem can be easily solved by the second strategy named overlapped mapping. Instead of mapping the grids onetoone, we extend the mapping into a one-to-many overlapped mapping. More specifically, a parameter called overlapped length (OL) is defined to denote the degree of overlapping. Now the new grid $G_{3i,j}$ is mapped to the old grids from $G_{i+x,j}$ to $G_{i+x+OL,j}$.

The most distinctive benefit of both direct and overlapped mapping strategies is that all shifting sources and targets can be easily computed without any search. However, the resulting mapping may not be optimal in terms of shifting cost since the same segment may be available on more than one SAP . The lowest cost shifting source may not be the one that is mapped directly. We therefore propose the third strategy named the shortest distance mapping strategy. After determining the segment assignment for the new zone, the distances between a new SAP in the zone and the $SAPs$ in the old zone with the assigned segments can be computed. We can therefore map the new SAP to the nearest neighbor in the old zone with the required segment. If there are more than one qualified neighboring $SAPs$ available to choose from, we follow the lowest shifted volume first strategy to balance the load.

c) Headlight Sharing

Headlight shifting only takes the headlight zone of one user into consideration. Segments not yet prefetched by any SAP of the old zone can only be

retrieved from the remote source. Since the same media may be viewed by more than one user at the same time, especially for hot medias, the headlight zones of different users may have many segments in common. If they overlap with each others, then it is very likely that we can find the needed segments from other zones with or without shifting. We call this idea headlight sharing since segments prefetched for a zone by an SAP are shared with neighboring SAPs with overlapping headlight zones. Once a requested segment is located in the neighborhood, the cost of prefetching from the remote server can also be saved.

To facilitate headlight sharing, a distributed index structure is constructed on each SAP for maintaining the availability of media segments on other SAPs. For ease of presentation, we assume that the set of all medias be M and the set of all segments of media i be G_i . A separate index table $T_i = \{S_i, M\}$ is maintained for each neighboring SAP from which an index message is received. S_i is the id of the SAP. $M = \{(mj, sk, t) | mj \in M, sk \in G_{mj}\}$ is the partial set of media and segments available on that SAP. Each segment has an expiration time t to indicate the period of validity. An SAP index keeps track of all the available index tables while a segment index maintains, for each segment, a list of SAPs from which the segment can be found. The information is obtained from the index exchange between SAPs. To avoid additional communication overhead, the index messages are piggybacked with the headlight prefetching messages. With the index ready, the segments prefetched by an SAP can be easily shared with other SAPs to facilitate headlight sharing.

The problem now comes to the efficient maintenance of the distributed index. To keep a complete index of all medias and segments on each SAP is clearly cost prohibitive. Therefore only selected information is exchanged based on the following strategies:

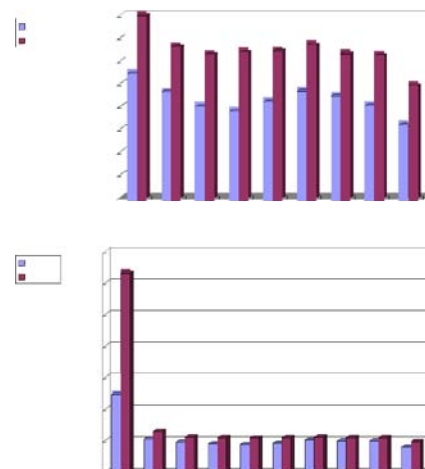
- *Random*: Randomly select part of the available segments to exchange index.
- *Local popularity*: Select those segments that are locally popular based on the rationale that they may be popular on other SAPs as well.
- *SAP specific popularity*: Segments that are popular among the users coming from the same SAP are selected to send the index to that particular SAP. This is to satisfy different request patterns of users from different cells.
- *SAP specific rareness*: Segments that are locally popular but rarely requested by the users coming from a particular SAP are selected to send the index to that SAP. This is because the segments available locally are to satisfy either the local requests or headlight prefetching. The availability of the prefetched segments are already known to the SAPs that send the requests. Therefore we only need to send the index of those segments that are not known to other SAPs.

Headlight sharing works closely with headlight prefetching and shifting. On receiving a segment request, be it from a local user or a prefetch request, an SAP looks up the index after a cache miss to see if the segment is available on other SAPs. The segment can then be retrieved from a nearby SAP rather than from the remote source. During the headlight shifting, those segments that are in the new assignment but no SAP to shift from can very likely be satisfied using the headlight sharing index. Since the index messages are piggybacked with the prefetching messages and the index search cost is quite low in comparison with segment transmission cost, the over head of headlight sharing is almost negligible. Later in Section 5, we will show that different combinations of these techniques result in significant performance improvement over on demand and simple look ahead prefetching in terms of response time, interruption rate and completion rate.

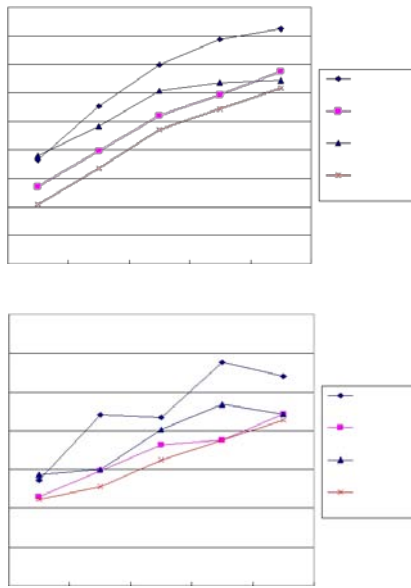
IV. SIMULATION AND EVALUATION

To evaluate the performance of our techniques, we have developed a Java based simulation environment. The set of simulation parameters and their value ranges are listed in Table 2.

The first set of experiments is to compare the performance of different combinations of headlight prefetching, shifting and sharing strategies. All strategies are tested under two request patterns: random and Zipf distribution. We are particular interested in the total number of interruptions during a media playback and the average download time of a media segment since they are the dominant factors affecting service quality experienced by the users. All strategies perform significantly better than traditional on demand strategy with simple lookahead prefetching. Since the inclusion of the later makes others hard to compare, we can see that shifting and sharing indeed improve the performance of headlight prefetching, especially when used together. All strategies achieve better performance under Zipf distribution since hot medias can be quickly shared with other SAPs.



Comparison of different headlight prefetching strategies on average segment download time.



Average segment download time with various numbers of medias.

To evaluate the scalability of our approach, we vary the number of total medias then measure the interruption and average segment download time. The combination of shifting and sharing, again leads to best result.

Traveling speed is an important factor in mobile environments since the faster the average speed, the less time we have to prepare media segments for users. When the average speed is below 40, both play interruption and average download time are relative independent of the variation in speed. When the average speed is too fast such that the prefetching can no longer catch up with the mobile users, we observe a significant increase in both play interruption and download time, especially the later.

The size of client cache also has significant impact on system performance since the larger the cache, the more we can share with others. We note that a cache size 1 does not necessarily imply zero interruption since a media is started after 2% buffering.

The saving is especially evident when the number of total medias is small since the chance of successful sharing increase. The set of headlight prefetching techniques are flexible and effective for media streaming services in mobile environments. Headlight prefetching provides a simple but uniform mechanism to allocate the resources on SAPs to the prefetching of needed media segments to the places where a user is most likely to be in the near future. When the movement pattern of a user does not follow the predicted track, headlight shifting and sharing leverage off the downloaded segments to quickly provide

seamless media streaming services along the new track. Simulation and performance evaluation demonstrate that the headlight prefetching technique with the help of both headlight shifting and sharing, provides the best performance overall.

V. CONCLUSIONS AND FUTURE WORK

We have proposed a new set of techniques to facilitate data management for media streaming in mobile environments. The headlight prefetching techniques provide good performance in comparison with traditional ondemand or simple prefetching techniques. To offer even better data management in response to the changes in user behavior such as access and moving patterns, we are developing an adaptive headlight prefetching technique such that the shape and size of the headlight zone can be dynamically adjusted to accommodate the changes in speed or direction. We are also developing a P2P dynamic chaining method for the sharing of information among peers to maximize cache utilization and streaming benefit.

REFERENCES

1. A. Boni, E. Launay, T. Mienville, and P. Stuckmann. Multimedia broadcast multicast service technology overview and service aspects. In 5th IEE International Conference on 3G Mobile Communication Technologies, pages 634–638, 2004.
2. Faria, J. A. Henriksson, E. Stare, and P. Talmola. DVBH: Digital broadcast services to handheld devices. Proceedings of the IEEE, 94(1):194–209, Jan 2006.
3. M. Hefeeda and B. K. Bhargava. On-demand media streaming over the internet. In 9th IEEE Workshop on Future Trends of Distributed Computing Systems (FTDCS 2003), pages 279–285, May 2003.
4. C.M. Huang and T. H. Hsu. A user aware prefetching mechanism for video streaming. World Wide Web Journal, 6(4):353–374, Dec 2003.
5. A. Kyriakidou, N. Karelou, and A. Delis. Video streaming for fast moving users in 3g mobile networks. In MobiDE 2005: 4th International ACM Workshop on Data Engineering for Wireless and Mobile Access, pages 65–72, 2005.
6. B. Li and K. H. Wang. NonStop: Continuous multimedia streaming in wireless ad hoc networks with node mobility.



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Implementation of Impulse Noise Reduction Method to Color Images using Fuzzy Logic

By G Venkateswara Rao, Satya P Kumar Somayajula, Dr. C.P.V.N.J Mohan Rao

Avanthi College of Engg & Tech, Tamaram

Abstract - Image Processing is a technique to enhance raw images received from cameras/sensors placed on satellites, space probes and aircrafts or pictures taken in normal day-to-day life for various applications. Impulse noise reduction method is one of the critical techniques to reduce the noise in color images. In this paper the impulse noise reduction method for color images by using Fuzzy Logic is implemented. Generally Grayscale algorithm is used to filter the impulse noise in corrupted color images by separate the each color component or using a vector-based approach where each pixel is considered as a single vector. In this paper the concepts of Fuzzy logic has been used in order to distinguish between noise and image characters and filter only the corrupted pixels while preserving the color and the edge sharpness. Due to this a good noise reduction performance is achieved. The main difference between this method and other classical noise reduction methods is that the color information is taken into account to develop a better impulse noise detection a noise reduction that filters only the corrupted pixels while preserving the color and the edge sharpness. The Fuzzy based impulse noise reduction method is implemented on set of selected images and the obtained results are presented.

Keywords : Image Processing, Impulse noise, Fuzzy logic.

GJCST Classification : I.4.3



Strictly as per the compliance and regulations of:



Implementation of Impulse Noise Reduction Method to Color Images using Fuzzy Logic

G Venkateswara Rao^a, Satya P Kumar Somayajula^a, Dr. C.P.V.N.J Mohan Rao^b

Abstract - Image Processing is a technique to enhance raw images received from cameras/sensors placed on satellites, space probes and aircrafts or pictures taken in normal day-to-day life for various applications. Impulse noise reduction method is one of the critical techniques to reduce the noise in color images. In this paper the impulse noise reduction method for color images by using Fuzzy Logic is implemented. Generally Grayscale algorithm is used to filter the impulse noise in corrupted color images by separate the each color component or using a vector-based approach where each pixel is considered as a single vector. In this paper the concepts of Fuzzy logic has been used in order to distinguish between noise and image characters and filter only the corrupted pixels while preserving the color and the edge sharpness. Due to this a good noise reduction performance is achieved. The main difference between this method and other classical noise reduction methods is that the color information is taken into account to develop a better impulse noise detection a noise reduction that filters only the corrupted pixels while preserving the color and the edge sharpness. The Fuzzy based impulse noise reduction method is implemented on set of selected images and the obtained results are presented.

Keywords : Image Processing, Impulse noise, Fuzzy logic.

I. INTRODUCTION

Processing of images which are digital in nature by a digital computer is called as digital image processing. Image Processing is a technique to enhance raw images received from cameras/sensors placed on satellites, space probes and aircrafts or pictures taken in normal day-to-day life for various applications. Various techniques have been developed in Image Processing during the last four to five decades. Most of the techniques are developed for enhancing images obtained from unmanned spacecrafts, space probes and military reconnaissance flights. Image Processing systems are becoming popular due to easy availability of powerful personnel computers, large size memory devices, graphics software etc. Image Processing is used in various applications such as remote sensing, medical imaging, film industry, military, etc.

Author^a : Asst. Professor, in MCA Department, Gayatri Vidya Parishad College for PG Courses, Rushikonda, Visakhapatnam, A.P., India.
E-mail : venkyintouch@gmail.com

Author^a : Asst. Professor, in CSE Department, Avanthi College of Engg & Tech, Tamaram, Visakhapatnam, A.P., India.
E-mail : balasriram1982@gmail.com

Author^b : Professor, in CSE Department, Principal of Avanthi Institute of Engineering & Technology, Narsipatnam

a) Color Models

The purpose of a color model is to facilitate the specification of colors in some standard, generally accepted way. In essence, a color model is a specification of a co-ordinate system and a subspace within that system where each color is represented by a single point.

b) Fuzzy Logic

In this paper Fuzzy logic concept has been used in order to distinguish between noise and image characters and filter only the corrupted pixels while preserving the color and the edge sharpness. Fuzzy set theory and fuzzy logic offer us powerful tools to represent and process human knowledge represented as fuzzy if-then rules. Fuzzy image processing has three main stages: 1) image fuzzification, 2) modification of membership values, and 3) image defuzzification. The fuzzification and defuzzification steps are due to the fact that we do not yet possess fuzzy hardware. Therefore, the coding of image data (fuzzification) and decoding of the results (defuzzification) are steps that make it possible to process images with fuzzy techniques. The main power of fuzzy image processing lies in the second step (modification of membership values).

II. METHOD

a) Implementing Filter to Remove Noise

This method consists of two phases viz., the Detection phase and De-noising phase. The result of the detection method is used to calculate the noise-free color component differences of each pixel. These differences are used by the noise reduction method so that the color component differences are preserved. We use the red-green-blue (RGB) color space as basic color space.

- 1) to the neighbors in the same color band and
- 2) to the color components of the two other color bands. If F_i denotes the input noisy image and O_i the original noise-free image at pixel position, then we can express the random-value impulse noise as

$$F_i^{col} = \begin{cases} O_i^{col}, & \text{with probability } 1 - \delta \\ \zeta_i^{col}, & \text{with probability } \delta \end{cases}$$

Where ζ_i^{col} is an identically distributed, independent random process with an arbitrary underlying probability density function. We consider the most used distribution: namely the uniform distribution,

where the noise was added to each color component independently. The indexes i and col indicate the 2-D pixel position and the color component, respectively, i.e., $col = R, G$ or B if the RGB-color space is used.

b) Impulse Noise Detection

1. Whether each color component value is similar to the neighbors in the same color band and
2. Whether the value differences in each color band corresponds to the value differences in the other bands. Since we are using the RGB color-space, the color of the image pixel at position i is denoted as the vector F_i which comprises its red (R), green (G), and blue (B) component, so $F_i = (F_i^R, F_i^G, F_i^B)$. Let us consider the use of a sliding filter window of size $n \times n$, with $n = 2c+1$ and $c \in \mathbb{N}$, which should be centered at the pixel under processing, denoted as F_o . For a 3×3 window, we will denote the neighboring pixels as F_1 to F_8 (i.e., from left to right and upper to lower corner). The color pixel under processing is always represented by $F_o = (F_o^R, F_o^G, F_o^B)$.

The Detection phase consists of the following seven steps

a) Calculation of absolute differential matrix

First, we compute the absolute value differences between the central pixel F_o and each color neighbor as follows:

$$\Delta F_k^R = |F_o^R - F_k^R|, \quad \Delta F_k^G = |F_o^G - F_k^G| \text{ and } \Delta F_k^B = |F_o^B - F_k^B|$$

where $k = 1, \dots, n^2-1$ and $\Delta F_k^R, \Delta F_k^G$ and ΔF_k^B denote the value difference with the color at position in the R, G, and B component, respectively.

b) Compute the fuzzy set $S1$ (membership degrees) for these differences

Now, we want to check if these differences can be considered as small. Since small is a linguistic term, it can be represented as a fuzzy set. Fuzzy sets, in turn, can be represented by a membership function. In order to compute the membership degree in the fuzzy set small we have to know the desired behavior, i.e., if the difference is relatively small then we want to have a large membership degree (the membership degree should decrease slowly), but after a certain point, we want to decrease the membership degree faster for each larger difference. Therefore, we have chosen the 1-S-membership function over other possible functions. This function is defined as follows:

$$1 - S(x) = \begin{cases} 1, & \text{if } x \leq \alpha_1 \\ 1 - 2 \left(\frac{x - \gamma_1}{\gamma_1 - \alpha_1} \right)^2, & \text{if } \alpha_1 < x \leq \frac{\alpha_1 + \gamma_1}{2} \\ 2 \left(\frac{x - \alpha_1}{\gamma_1 - \alpha_1} \right)^2, & \text{if } \frac{\alpha_1 + \gamma_1}{2} < x \leq \gamma_1 \\ 0, & \text{if } x > \gamma_1 \end{cases}$$

where it has been experimentally found that $\alpha_1=10$ and $\gamma_1=70$ receive satisfying results in terms of noise detection. In this case, we denote 1-S by S_1 , so that $S_1(\Delta F_k^R)$, $S_1(\Delta F_k^G)$, $S_1(\Delta F_k^B)$ denote the membership degrees in the fuzzy set $small_1$ of the computed differences with respect to the color at position k .

c) Calculate the degree of similarity μ^R, μ^G, μ^B

Now, we use the values $S_1(\Delta F_k^R)$, $S_1(\Delta F_k^G)$, $S_1(\Delta F_k^B)$ for $k = 1, \dots, n^2-1$ to decide whether the values F_o^R , F_o^G and F_o^B are similar to its component neighbors. The number k of considered neighbors will be a parameter of the filter. So, we apply a fuzzy conjunction operator (fuzzy AND operation represented here by the triangular product t-norm among the first k ordered membership degrees in the fuzzy set $small_1$. The conjunction is calculated as follows:

$$\mu^R = \prod_{j=1}^K S_1(\Delta F_{(j)}^R)$$

where μ^R denotes the degree of similarity between F_o^R and K the n -nearest neighbors.

d) From $S1$ calculate $S1(RG, GB, BR)$ i.e differences among R, G, B components

Besides the first step of the detection method, i.e., checking if the central pixel is similar to its local neighborhood or not, we investigate whether the color components are correlated which each other or not. In other words, we determine whether the local differences in the R component neighborhood corresponds to the differences in the G and B component. we compute the absolute value of the difference between the membership degrees in the fuzzy set $small_1$ for the red and the green and for the red and the blue components, i.e., $|S_1(\Delta F_k^R) - S_1(\Delta F_k^G)|$ and $|S_1(\Delta F_k^R) - S_1(\Delta F_k^B)|$ where $k = 1, \dots, n^2-1$, respectively.

e) Compute the fuzzy set $S2$ (membership degrees) for these differences

Now, in order to see if the computed differences are small we compute their fuzzy membership degrees in the fuzzy set $small_2$. The 1-S membership function is also used but now we used $\alpha_2=0.01$ and $\gamma_2=0.15$ and, which also have been determined experimentally. In this case we denote the membership function as S_2

f) Calculate the joint similarity $\mu^{RG}, \mu^{RB}, \mu^{BG}$

we calculate

$$\mu_k^{RG} = S_2(|S_1(\Delta F_k^R) - S_1(\Delta F_k^G)|)$$

$$\mu_k^{RB} = S_2(|S_1(\Delta F_k^R) - S_1(\Delta F_k^B)|)$$

where μ_k^{RG} and μ_k^{RB} denote the degree in which the local difference (between the center pixel and the pixel at position j) in the red component is similar to the local difference in the green and blue components. The obtained degrees μ_k^{RG} and μ_k^{RB} are sorted again sorted in descending order, where $\mu_{(j)}^{RG}$ and $\mu_{(j)}^{RB}$ denote the values ranked at the k^{th} position. Consequently, the *joint similarity* with respect to k neighbors is computed as

$$\mu^{RG} = \prod_{j=1}^K \mu_{(j)}^{RG}, \quad \mu^{RB} = \prod_{j=1}^K \mu_{(j)}^{RB}$$

where μ^{RG} and μ^{GB} denote the degree in which the local differences for the red component are similar to the local differences in the green and blue components, respectively. Notice that if F_0^R is noisy and F_0^G and F_0^B are noise-free, then the local differences can hardly be similar, and, therefore, low values of μ^{RG} and μ^{GB} are expected.

g) *Calculation of Noise-Free degree*

$$NF_{F_0^R}, NF_{F_0^G}, NF_{F_0^B}$$

Finally, the membership degree in the fuzzy set *noise-free* for F_0^R is computed using the following fuzzy rule

Fuzzy Rule 1: Defining the membership degrees $NF_{F_0^R}$ for the red component F_0^R in the fuzzy set *noise-free*

IF μ^R is large AND μ^{RG} large AND μ^G is large
OR

μ^R is large AND μ^{RB} large AND μ^B is large

THEN

the noise – free degree F_0^R is large

A color component is considered as noise-free if

- 1) it is similar to some of its neighbor values (μ^R) and
- 2) the local differences with respect to some of its neighbors are similar to the local differences in some of the other color components (μ^{RG} and μ^{GB}).

In fuzzy logic, triangular norms and co-norms are used to represent conjunctions and disjunctions respectively. Since we use the product triangular norm to represent the fuzzy AND (conjunction) operator and the probabilistic sum co-norm to represent the fuzzy OR (disjunction) operator the noise-free degree of F_0^R which we denote as $NF_{F_0^R}$ is computed as follows

$$NF_{F_0^R} = \mu^R \mu^{RG} \mu^G + \mu^R \mu^{RB} \mu^B - \mu^R \mu^{RG} \mu^G \mu^{RB} \mu^B$$

Analogously to the calculation of noise-free degree for the red component described above, we obtain the noise-free degrees of F_0^G and F_0^B denoted as $NF_{F_0^G}$ and $NF_{F_0^B}$ as follows

$$NF_{F_0^G} = \mu^G \mu^{RG} \mu^R + \mu^G \mu^{GB} \mu^B - \mu^G \mu^{RG} \mu^R \mu^{GB} \mu^B$$

$$NF_{F_0^B} = \mu^B \mu^{RB} \mu^R + \mu^B \mu^{GB} \mu^G - \mu^B \mu^{RB} \mu^R \mu^{GB} \mu^G$$

In fuzzy logic, involutive negators are commonly used to represent negations. We use the standard negator $N_s(x) = 1 - x$, with $x \in [0,1]$. By using this negation, we can also derive the membership degree in the fuzzy set *noise* for each color component, i.e., $NF_0^R = 1 - NF_{F_0^R}$, where $NF_{F_0^R}$ denotes the membership degree in the fuzzy set *noise*.

Algorithm for Impulse Noise Generator

Step1: Read the pixels from image ,we take some temporary variable initialize to zero.

Step2: For Red

Step 2.1: check the condition if temporary variable equal to zero assign color code 0x00ff0000.

Step 2.2: check the condition if temporary variable equal to one assign color code 0xff00ffff.

Step 2.3: repeat Step 2.1 and Step 2.2 until red pixels are encountered.

Step3: For Green

Step 3.1: check the condition if temporary variable equal to zero assign color code 0x0000ff00

Step 3.2: check the condition if temporary variable equal to one assign color code 0xffff00ff

Step 3.3: repeat Step 3.1 and Step 3.2 until green pixels are encountered .

Step4: For Blue

Step 4.1: check the condition if temporary variable equal to zero assign color code 0x000000ff

Step 4.2: check the condition if temporary variable equal to one assign color code 0xfffffff0

Step 4.3: repeat Step 4.1 and Step 4.2 until blue pixels are encountered.

III. RESULTS ANALYSIS

Different images as inputs are taken and apply this algorithm on these images and obtained the PSNR values .All these values are tabulated in table:1.

	
Original Mushroom image	Image corrupted with 5% random-value impulse noise. PSNR = 15
	
output of the filter with a 3x3 window and K = 2; PSNR = 28	the output of the proposed two-step filter PSNR = 30

Figure 1 : Noise Detection For Random-Value Impulse Noise

	
Original Mushroom image	Image corrupted with 5% fixed-value impulse noise. PSNR =22
	
output of the filter with a 3x3 window and K = 2; PSNR = 38	output of the proposed two-step filter. PSNR = 35

Figure 2 : Denoising - 5 % Fixed-Value Impulse Noise

PSNR VALUES OF A COLOR IMAGE			
RIN %	Filter	FIN %	Filter
1	36.089537	1	30.2734961
2	33.6586222	2	28.1291336
3	32.1085337	3	26.5705585
4	30.6483222	4	25.5990663
5	30.136797	5	24.6834733
6	28.0413943	6	23.9445168
7	27.5204845	7	23.1597035
8	26.4836001	8	22.9446623
9	25.7170883	9	22.6951294
10	25.2244423	10	22.1748394

Table 1 : PSNR Valued of lena image corrupted with (RIN & FIN both ranging from 1 to 10)

IV. CONCLUSION

In this paper, a new fuzzy filter for impulse noise reduction in color images is presented. The main difference between the proposed method (denoted as INR) and other classical noise reduction method is that the color information is taken into account in a more appropriate way. This method also illustrates that color images should be treated differently than grayscale images in order to increase the visual performance.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Stefan Schulte, Samuel Morillas, Valentín Gregori, and Etienne E. Kerre "A New Fuzzy Color Correlated Impulse Noise Reduction Method" IEEE transactions on image processing, vol.16, no.10, oct 2007.
2. S. Schulte, V. De Witte, M. Nachtegaal, D. Van der Weken, and E. E.Kerre, "Fuzzy random impulse noise reduction method," Fuzzy Sets Syst., vol. 158,

- no. 3, pp. 270–283, Jan. 2007.
3. S. Schulte, M. Nachtegael, V. De Witte, D. Van der Weken, and E. E.Kerre, "A fuzzy impulse noise detection and reduction method," *IEEE Trans. Image Process.*, vol. 15, no. 5, pp. 1153–1162, May 2006.
4. Z. H. Ma, H. R. Wu, and B. Qiu, "A robust structure-adaptive hybrid vector filter for color image restoration," *IEEE Trans. Image Process.*, vol. 14, no. 12, pp. 1990–2001, Dec. 2005.
5. H. Xu, G. Zhu, H. Peng, and D.Wang, "Adaptive fuzzy switching filter for images corrupted by impulse noise," *Pattern Recognit. Lett.*, vol. 25, pp. 1657–1663, Nov. 2004.
6. R. Lukac, K. N. Plataniotis, B. Smolka, and A. N. Venetsanopoulos, "A multichannel order-statistic technique for cDNA microarray image processing," *IEEE Trans. Nanobiosci.*, vol. 3, no. 4, pp. 272–285, Dec. 2004.
7. R. Lukac, K. N. Plataniotis, B. Smolka, and A. N. Venetsanopoulos, "Generalized selection weighted vector filters," *EURASIP J. Appl. Signal Process.*, vol. 2004, no. 12, pp. 1870–1885, Sept. 2004.
8. J. H. Wang, W. J. Liu, and L. D. Lin, "Histogram-Based fuzzy filter for image restoration," *IEEE Trans. Syst., Man, Cybern. B, Cybern.*, vol. 32, no. 2, pp. 230–238, Apr. 2002.
9. Chatzis and I. Pitas, "Fuzzy scalar and vector median filters based on fuzzy distances," *IEEE Trans. Image Process.*, vol. 8, no. 5, pp. 731–734, May 1999.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 22 Version 1.0 December 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Classification of Aerial Data Based on Data Mining Clustering Algorithm

By Prof.G.Ramaswamy, Dr. Vuda.Sreenivasarao, Dr.Popuri.Ramesh Babu,
P.V.S.S.Gangadhar

Defence University college

Abstract - The Aerial data contains data periodically observed with parameters of texture (min, max), flora, and density (min, max). The proposed Aerial prediction system cluster and analyze, three input features that is average texture, flora, average density according to number of days to predict Aerial for Surveillance applications. The proposed system realizes the k-means clustering algorithm for grouping similar features based on user intended period, further the system analyze using PCA (Principal Component Analysis) on same data.

Keywords : *Data mining, Aerial data, cluster algorithm, Principal Component Analysis.*

GJCST Classification : *H.2.8*



Strictly as per the compliance and regulations of:



A Classification of Arial Data Based on Data Mining Clustering Algorithm

Prof.G.Ramaswamy^α, Dr. Vuda.Sreenivasarao^Ω, Dr.Popuri.Ramesh Babu^β, P.V.S.S.Gangadhar^ψ

Abstract -The Arial data contains date periodically observed with parameters of texture (min, max), flora, and density (min, max). The proposed Arial prediction system cluster and analyze, three input features that is average texture, flora, average density according to number of days to predict Arial for Surveillance applications. The proposed system realizes the k-means clustering algorithm for grouping similar features based on user intended period, further the system analyze using PCA (Principal Component Analysis) on same data.

Keywords : Data mining, Arial data, cluster algorithm, Principal Component Analysis.

I. INTRODUCTION

The Arial data contains date periodically observed with parameters of texture (min, max), flora, and density (min, max). The amount of data kept in computer files and databases is growing at a phenomenal rate. At the same time, the users of these data are expecting more sophisticated information from them. For example a marketing manager is no longer satisfied with a simple listing of marketing contacts, but wants detailed information about customers past purchases as well as predictions of future purchases. Simple structured/query language queries are not adequate to support these increased demands for information. Data mining steps in to solve these needs. Data mining is often defined as finding hidden information in a database. Data mining involves the use of sophisticated data analysis tools to discover previously unknown, valid patterns and relationships in large data sets. These tools can include statistical models, mathematical algorithms, and machine learning methods. Consequently, data mining consists of more than collecting and managing data, it also includes analysis and prediction. Data mining can be performed on data represented in quantitative, textual, or multimedia forms. Data mining applications can use a variety of parameters to examine the data. They include association, sequence or path analysis, classification, clustering, and forecasting. This paper deals with implementation of an automated Arial prediction system for agriculture applications using data mining tools such as clustering (k-means algorithm) and Principle Component Analysis (PCA). For making accurate

decision on large observation is an important factor, but with increasing information the clustering algorithm faces various limitations/problems. Among them current clustering techniques do not address all the requirements adequately (and concurrently). Dealing with large number of dimensions and large number of data items can be problematic because of time complexity. The effectiveness of the method depends on the definition of "distance" (for distance-based clustering). If an obvious distance measure doesn't exist we must "define" it, which is not always easy, especially in multi-dimensional spaces. The result of the clustering algorithm (that in many cases can be arbitrary itself) can be interpreted in different ways. The main objective of this paper is to develop an accurate and efficient Arial prediction system for agriculture applications using data mining tools such as clustering (k-means algorithm) and PCA.

II. DATA CLUSTERING

Clustering is a divided number of groups of similar data objects. Each group called cluster, consists of objects that are similar between themselves and dissimilar to objects of other groups. Representing the data by fewer clusters necessarily loses certain fine details, but achieves simplification. It models data by its clusters. Data modeling puts clustering in a historical perspective rooted in numerical analysis, mathematics and statistics. From a machine learning perspective clusters correspond to hidden patterns, the search for clusters is unsupervised learning, and the resulting system represents a data concept. In practical perspective clustering plays an outstanding performance in data mining applications such as, computational biology ,information retrieval and text mining, scientific data exploration ,marketing, medical diagnostics, spatial database applications, and Web analysis, etc.

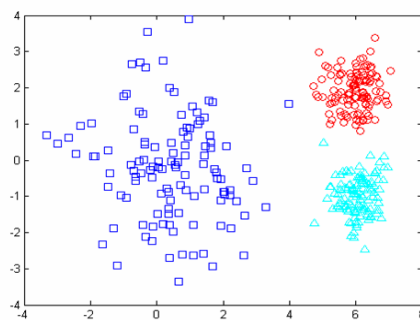


Figure1 : Clusters distribution of a data set

Author^α : Principal, Priyadarshini Inst.Of Tech.&Sci, Guntur, India.

Author^Ω : Professor in CIT, Defence University college, Debrezeit, Ethiopia.

Author^β : Dean Faculty of Engineering, Malineni Perumallu Education, Society's Group of Institutions, India.

Author^ψ : Scientist C, NIC, Govt of India, India.

Clustering is the subject of active research in several fields such as statistics, pattern recognition, data mining, grouping and decision making, pattern classification, bio informatics and machine learning. A very important characteristic of most of these application domains is that the size of the data involved is very large. So, clustering algorithms used in these application areas should be able to handle large data sets of sizes ranging from gigabytes to terabytes and even pica bytes. Typically, clustering algorithms paper on pattern matrices, where each row of the matrix corresponds to a distinct pattern and each column corresponds to a feature. Most of the early paper on clustering dealt with the problem of grouping small data sets, where the benchmark data sets used to demonstrate the performance of the clustering algorithms were having a few hundreds of patterns and a few tens of features. Fisher's Iris data is one of the most frequently used benchmark data sets. This data has three classes, where each class has 50 patterns and each pattern is represented using four features. However, several real-world problems of current interest are very large in terms of the pattern matrices involved. For example, in data mining and web mining the number

of patterns is typically very large, whereas in clustering biological sequences, the number of features involved is very large. Cluster analysis is a way to examine similarities and dissimilarities of observations or objects. Data often fall naturally into groups, or clusters, of observations, where the characteristics of objects in the same cluster are similar and the characteristics of objects in different clusters are dissimilar. In one of the earliest books on data clustering, Underberg defines cluster analysis as a task, which aims to finding of natural groups from a data set, when little or nothing is known about the category structure. Bailey, who surveys the methodology from the sociological perspective, defines that cluster analysis seeks to divide a set of objects into a small number of relatively homogeneous groups on the basis of their similarity over N variables. N is the total number of variables in this case.

III. SYSTEM ARCHITECTURE

The Arial prediction system architecture is shown in figure 2. The overall system design consists of Input (Arial Data), modified input, Feature selection, Clustering Using k-means on Selected Feature And PCA on Selected Feature.

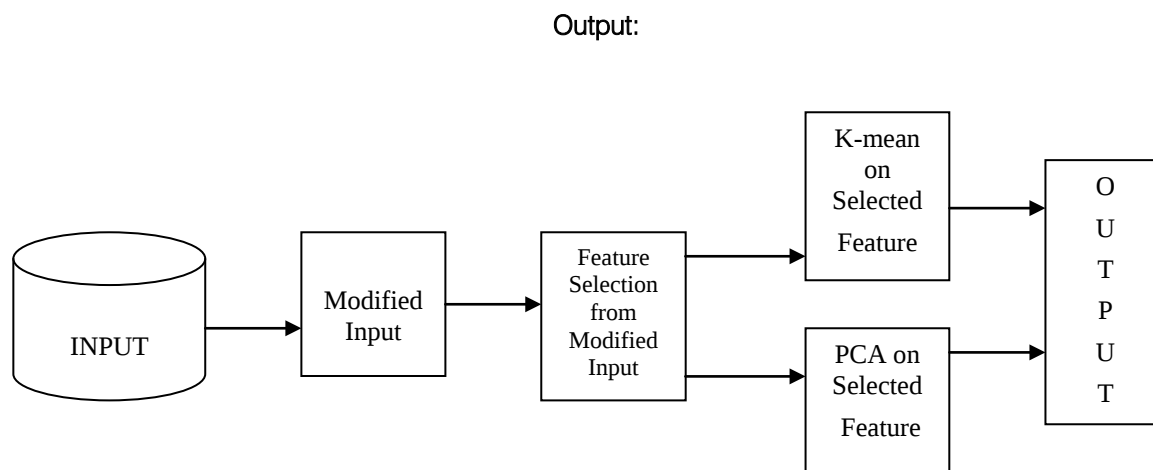


Figure 2 : Architecture of Arial prediction system

Apply k-means clustering algorithm on selected feature. Clustering is a division of data into groups of similar objects. Each group called cluster, consists of objects that are similar between themselves and dissimilar to objects of other groups. K-means is a typical unsupervised learning clustering algorithm. It partitions a set of data into k clusters. However, it assumes that k is known in advance. Following is the summary of the algorithm:

1. Put K points into the representation of space by the objects that are being clustered. These points represent group Centroids.
2. Allocated each object to the group that has the closest centroids
3. When all objects have been allocated, again calculation of the positions of the K centroids.
4. Repeating of Steps 2 and 3 up to the centroids no longer move. This produces a separation of the objects into number of groups from which the metric to be minimized can be calculated.

The papering flow of "clustering using k-means on selected data" is explained by flowchart shown in figure 3. Start the procedure, next accept starting period, ending period data from user, and accept k value from user. Accept clustering data option,

- 4) Average Texture
- 5) Flora
- 6) Average Density

Cluster Arial data using k-means algorithm on selected feature (between starting and ending period days). If select 4 as our option than average Texture data is cluster, if select 5 than rain, if select 6 than average density data is cluster using k-means clustering algorithm. Display final results and Stop the procedure.

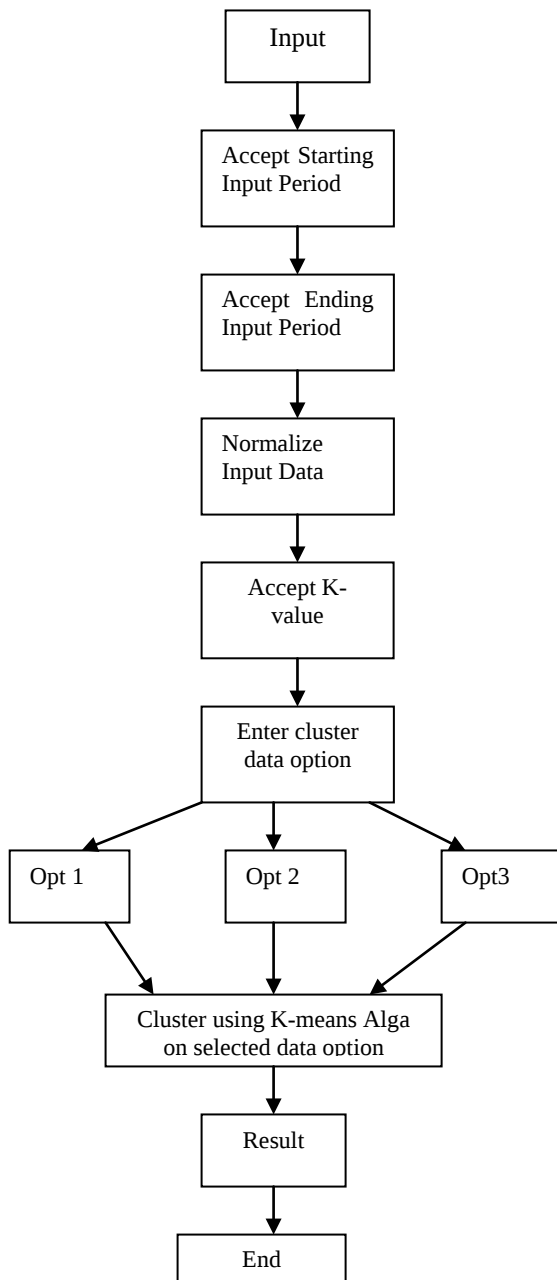


Figure 3 : Clustering Arial Data using K-means Algorithm

Apply PCA algorithm on selected feature. The papering flow of “PCA on selected data” is explained by flowchart shown in figure 4.

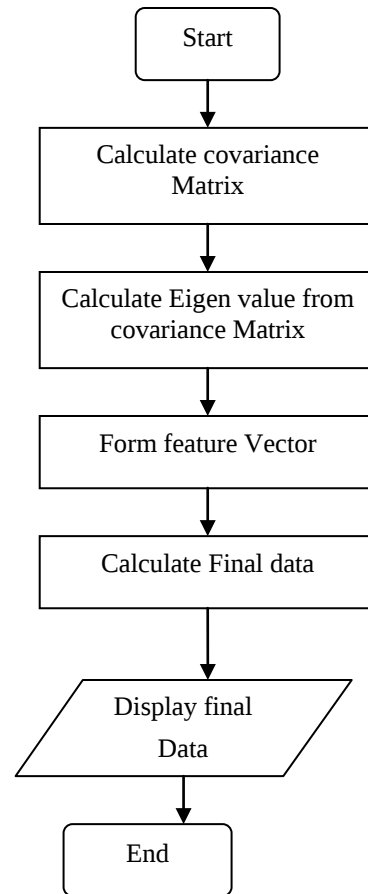


Figure 4 : PCA method on Arial data

Start the procedure next accept starting period, ending period data from user, and accept k value from user. Accept PCA data option, Perform PCA Method on selected data option. Subtract mean from normalize option data store in adjusted data variable. Calculate adjusted data covariance matrix. Calculate Eigen vector, value from covariance matrix, next select eigenvector with highest Eigen value is feature vector. Next calculate final data, display final data and stop the procedure.

IV. RESULT ANALYSIS

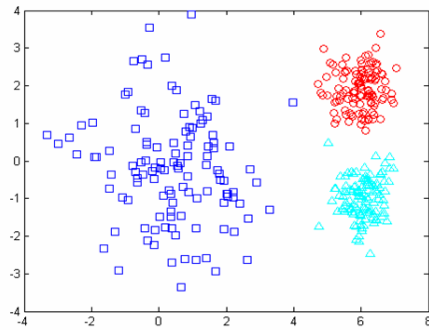


Figure 5 : distributed data set of the read information

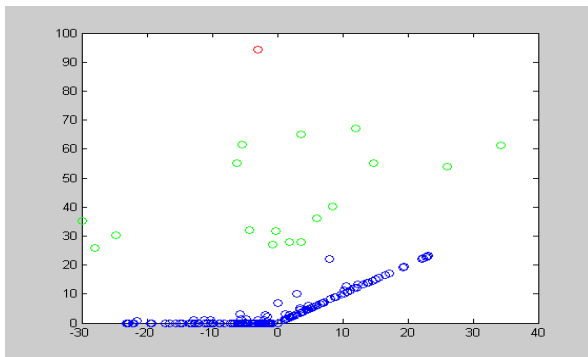


Figure 6 : the classified Component analysis data set for the reference model

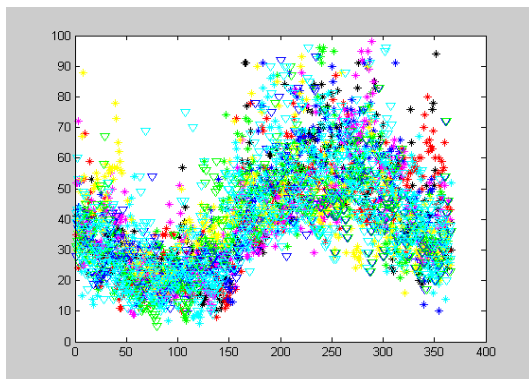


Figure 7 : obtained distribution reference for the given input data

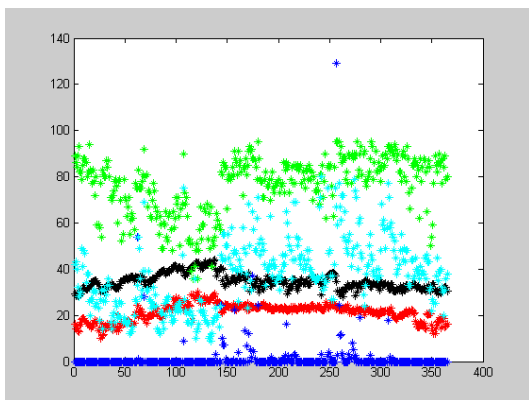


Figure 8 : Observations for obtained classified data set for the input information

V. CONCLUSION

The Arial data contains date periodically observed with parameters of texture (min, max), flora, and density (min, max).Farmer needs timely and accurate Arial data. In order to achieve this, data should be continuously recorded from stations that are properly identified, manned by trained staff or automated with regular maintenance, in good papering order and secure from tampering. The stations should also have a long history and not be prone to relocation. The collection and archiving of Arial data is important because it provides an economic benefit but the local/national economic needs are not as dependent on high data quality as is the Arial risk market. In this study, it was found that the data mining tools could enable experts to predict Arial with satisfying accuracy using as input the Arial parameters of the previous years. The K-means clustering and PCA algorithms are suggested and tested for period of 11 years with multiple features to early prediction of Arial for agriculture applications.

REFERENCES REFERENCES REFERENCIAS

1. K. Jain and R. C. Dubes, Algorithms for Clustering Data. Prentice-Hall, New Jersey, 1988.
2. R. A. Fisher. The Use of Multiple Measurements on Taxonomic Problems, Annals of Eugenics, vol. 7, 179-188, 1936.
3. David J. Hand, Heikki Mannila and Padhraic Smyth. Principles of Data Mining, MIT Press, 2001.
4. R. W. Cooley. Web usage mining: discovery and application of interesting patterns from web data. PhD thesis, University of Minnesota, USA, 2000.
5. M. R. Anderberg, Cluster analysis for applications, Academic Press, Inc., London, 1973.
6. K. D. Bailey, Cluster analysis, Sociological Methodology, (1975), 59-128.
7. R. J. Little and D. B. Rubin, Statistical analysis with missing data, John Wiley & Sons, 1987.
8. Guha. S, Rastogi. R, and Shim, 1998. Cure: An efficient clustering algorithm for large databases. In Proceedings of the ACM Sigmod Conference, 73-84, Seattle, WA.
9. Berry. M. W and Browne. M, 1999, Understanding Search Engines: Mathematical Modeling and Text Retrieval. SIAM.
10. P. Bradley, C. Reina, and U. Fayyad, Clustering very large databases using EM mixture models, in Proceedings of 15th International Conference on Pattern Recognition (ICPR'00), vol. 2, 2000, 76-80.
11. J. Macqueen, Some methods for classification and analysis of multivariate observations, in Proc. of the Fifth Berkeley Symposium on Mathematical Statistics and Probability, 1967, 281-297. Lindsay I Smith, "A tutorial on Principal Components Analysis", February 26, 2002.
12. M.C. Todd, C.Kidd, D. Kniveton and T. J. Bellerby, A combined satellite infrared and passive microwave

- technique for estimation of small-scale flora, J.Atmos, Oceanic technol., vol. 18, 724-75, 2001.
13. R.F. Adler and A.J. Negri, A satellite technique to estimate tropical convective and stratiform flora, J. Appl. Meteorol., vol. 27, 30-51, 1988.
 14. D.I.F Grimes, E. Coppola, M.Verdecchia, and G. Visconti, A neural netpaper approach to real-time flora estimation for Africa using satellite data, J. Hydrometeorol., vol. 4, 1119-1133, 2003.
 15. Degaetano, A spatial grouping of United States climate stations using a hybrid clustering approach. International Journal of Climatology, Chichester, Vol.21, 791-807, 2001.



GLOBAL JOURNALS INC. (US) GUIDELINES HANDBOOK 2011

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 GJMBR 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 brevity.

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.



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 a 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.



INDEX

A

Adaboost · 98, 100
Approximate · 6
Associative · 51, 53, 55, 56, 59, 61
attributes · 43, 44, 46, 47, 53, 55, 56, 57
Augmented · 2, 97
Automatic · 2, 4, 14, 15

B

behavioral · 28, 30
biological · 65, 89, 127
biorthogonal · 22
bottleneck · 24, 104
boundary · 28, 55, 56, 65, 80

C

Calculi · 2, 73, 81, 85, 86, 87
clinicians · 75
cluster · 79, 125, 127, 129
coefficients · 18, 19, 20, 22, 23, 24
coherence · 18, 19
Collaborative · 2, 97
complications · 48
component · 8, 18, 57, 67, 78, 101, 116, 118, 120, 121
computational · 4, 5, 18, 24, 73, 84, 125
conformance · 31, 35, 99
crosscutting · 33
cumbersome · 97
cylindrical · 18, 22, 23, 24

D

decompose · 24
defuzzification · 67, 116
density · 117, 125, 129, 130
Descriptive · 41, 86
diagnosis · 75, 76
Dialects · 16
diminution · 41, 42, 46, 47, 48

E

enhancement · 6, 77, 79, 99
enriches · 95

epidemics · 55
Equalization · 73, 77, 81
Estimation · 2, 63, 69, 70
execution · 29, 31, 36, 37, 38, 55
expressiveness · 28, 31, 34, 35, 38
Extraction · 43, 44

F

fingertips · 101
Framework · 49, 61, 89

G

Gesture · 97, 101
gradient · 100

H

Headlight · 2, 104, 107, 108, 110, 111, 113, 114
healthcare · 2, 51
heuristics · 9, 14
histogram · 76, 79, 81, 100
Homogeneous · 76

I

illuminance · 104, 107, 108, 109, 110
imprecision · 67
improvisational · 8
Impulse · 2, 116, 118, 121, 122
Indictor · 2, 73, 77
indistinguishable · 5
inspiration · 65
instantiations · 30, 33, 34
intangibly · 41
Intelligence · 2, 15, 59, 63, 65, 66, 101
interactions · 29, 92

J

juxtaposed · 69

L

luminance · 99, 100

M

mandatory · 31, 35, 36
matrix · 11, 12, 14, 22, 23, 24, 118, 127, 129
mechanism · 28, 73, 104, 108, 114, 115
merged · 31, 47
morphological · 76, 79, 81
multiresolution · 18, 26
multivariate · 130

N

neighborhood · 112, 118
nonlinear · 69, 100

O

ontology · 92
Organizations · 2, 89, 95
Oriented · 28, 30, 39, 49
overlapped · 111

P

parameters · 8, 14, 63, 66, 67, 69, 73, 75, 76, 77, 80, 81, 93, 109, 113, 125, 130
Participants · 93, 95
photography · 20
piggybacked · 112
Plenoptic · 26
Preassigned · 49
precomputed · 22, 23
Predictive · 2, 41, 51, 61
preparatory · 99
professional · 75
protocol · 36, 37, 38

Q

quantitative · 55, 56, 125

R

rearranged · 44, 46, 47
Recognition · 4, 14, 15, 16, 48, 101, 130
reconnaissance · 116
Recurrent · 41, 44

S

satellites · 116
scalability · 28, 29, 31, 33, 34, 35, 38, 114
segments · 42, 104, 106, 107, 108, 109, 110, 111, 112, 113, 114
Semidetached · 63
sensitivity · 84, 85
Serializaibility · 37
strategies · 6, 14, 91, 93, 95, 104, 106, 107, 111, 112, 113, 114
Swarm · 2, 63, 65, 66, 67, 69, 70
Symposium · 15, 49, 59, 61, 85, 86, 87, 130

T

Techniques · 33, 66, 70, 85
transcriptor · 4, 6, 14
Transformation · 2, 28, 33, 34, 39
traversals · 79

U

Ultrasound · 2, 73, 74, 85, 86
understandable · 7, 28, 89, 93
unsupervised · 125, 127
urinary · 73, 75

V

variables · 53, 127

W

wavelet · 18, 19, 20, 22, 23, 24, 25, 26, 86
Weighted · 41, 42, 44, 47, 48, 49, 56, 59, 61

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





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

© 2011 by Global Journals