

GLOBAL JOURNAL

OF COMPUTER SCIENCE AND TECHNOLOGY

DISCOVERING THOUGHTS AND INVENTING FUTURE

Technology
Reforming
Ideas

September 2011

Pinnacles

Speech Recognition System

Wireless Sensor Network

Vehicle Routing Problem

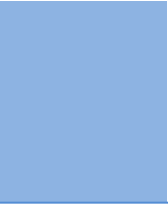
Service-Oriented Architecture

The Volume 11

Issue 15
VERSION 1.0



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY

VOLUME 11 ISSUE 15 (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://www.globaljournals.org/global-journals-research-portal/guideline/terms-and-conditions/menu-id-260/>.

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

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

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

Global Journals Inc.

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

Sponsors: Open Association of Research Society

Open Scientific Standards

Publisher's Headquarters office

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

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

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

Offset Typesetting

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

Packaging & Continental Dispatching

Global Journals, India

Find a correspondence nodal officer near you

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

eContacts

Press Inquiries: press@globaljournals.org

Investor Inquiries: investers@globaljournals.org

Technical Support: technology@globaljournals.org

Media & Releases: media@globaljournals.org

Pricing (Including by Air Parcel Charges):

For Authors:

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

Yearly Subscription (Personal & Institutional):

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

EDITORIAL BOARD MEMBERS (HON.)

John A. Hamilton,"Drew" Jr.,

Ph.D., Professor, Management
Computer Science and Software
Engineering
Director, Information Assurance
Laboratory
Auburn University

Dr. Henry Hexmoor

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

Dr. Osman Balci, Professor

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

Yogita Bajpai

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

Dr. T. David A. Forbes

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

Dr. Wenying Feng

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

Dr. Thomas Wischgoll

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

Dr. Abdurrahman Arslanyilmaz

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

Dr. Xiaohong He

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

Burcin Becerik-Gerber

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

Dr. Bart Lambrecht

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

Dr. Carlos García Pont

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

Dr. Fotini Labropulu

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

Dr. Lynn Lim

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

Dr. Mihaly Mezei

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

Dr. Söhnke M. Bartram

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

Dr. Miguel Angel Ariño

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

Philip G. Moscoso

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

Dr. Sanjay Dixit, M.D.

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

Dr. Han-Xiang Deng

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

Dr. Pina C. Sanelli

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

Dr. Roberto Sanchez

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

Dr. Wen-Yih Sun

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

Dr. Michael R. Rudnick

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

Dr. Bassey Benjamin Esu

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

Dr. Aziz M. Barbar, Ph.D.

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

PRESIDENT EDITOR (HON.)

Dr. George Perry, (Neuroscientist)

Dean and Professor, College of Sciences

Denham Harman Research Award (American Aging Association)

ISI Highly Cited Researcher, Iberoamerican Molecular Biology Organization

AAAS Fellow, Correspondent Member of Spanish Royal Academy of Sciences

University of Texas at San Antonio

Postdoctoral Fellow (Department of Cell Biology)

Baylor College of Medicine

Houston, Texas, United States

CHIEF AUTHOR (HON.)

Dr. R.K. Dixit

M.Sc., Ph.D., FICCT

Chief Author, India

Email: authorind@computerresearch.org

DEAN & EDITOR-IN-CHIEF (HON.)

Vivek Dubey(HON.)

MS (Industrial Engineering),

MS (Mechanical Engineering)

University of Wisconsin, FICCT

Editor-in-Chief, USA

editorusa@computerresearch.org

Sangita Dixit

M.Sc., FICCT

Dean & Chancellor (Asia Pacific)

deanind@computerresearch.org

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. Speech Recognition System Based on Hidden Markov Model Concerning the Moroccan Dialect DARIJA. *1-5*
 - 2. A Fine Grained Access Control Model Based on Diverse Attributes. *7-12*
 - 3. Recognition of Handwritten Tifinagh Characters Using a Multilayer Neural Networks and Hidden Markov Model. *13-19*
 - 4. Ubiquitous Life Care Integrates Wireless Sensor Network And Cloud Computing With Security. *21-27*
 - 5. Recognition of Tifinagh Characters Using Self Organizing Map And Fuzzy K-Nearest Neighbor. *29-33*
 - 6. A Framework for Costing Service-Oriented Architecture (SOA) Projects Using Work Breakdown Structure (WBS) Approach. *35-47*
 - 7. A Novel Approach for Always Best Connected in Future Wireless Networks. *49-53*
 - 8. Improved Energy and Latency Efficient MAC Scheme for Dense Wireless Sensor Networks. *55-59*
 - 9. Optimal High Performance Self Cascode CMOS Current Mirror. *61-63*
 - 10. Towards The Solution of Variants of Vehicle Routing Problem. *65-72*
 - 11. Aavoiding Loops and Packet Losses in ISP Networks. *73-77*
-
- 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 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Speech Recognition System Based On Hidden Markov Model Concerning the Moroccan Dialect DARIJA

By A. El Ghazi, C. Daoui, N. Idrissi, M. Fakir , B. Bouikhalene

Abstract - In this work, we present a system for automatic speech recognition on the Moroccan dialect. We used the hidden Markov model to model the phonetic units corresponding to words taken from the training base. The results obtained are very encouraging given the size of the training set and the number of people taken to the registration. To demonstrate the flexibility of the hidden Markov model we conducted a comparison of results obtained by the latter and dynamic programming.

Keywords : Hidden Markov Model (HMM), MFCC, DTW, Acoustic vectors .

GJCST Classification : I.2.7, G.3



Strictly as per the compliance and regulations of:



Speech Recognition System Based On Hidden Markov Model Concerning the Moroccan Dialect DARIJA

A. El Ghazi^a, C. Daoui^Q, N. Idrissi^β, M. Fakir ^ψ, B. Bouikhalene[¥]

Abstract - In this work, we present a system for automatic speech recognition on the Moroccan dialect. We used the hidden Markov model to model the phonetic units corresponding to words taken from the training base. The results obtained are very encouraging given the size of the training set and the number of people taken to the registration. To demonstrate the flexibility of the hidden Markov model we conducted a comparison of results obtained by the latter and dynamic programming.

Keywords : Hidden Markov Model (HMM), MFCC, DTW, Acoustic vectors.

I. INTRODUCTION

The system of automatic speech recognition (ASR) can transcribe a voice message, extracting linguistic information from an audio signal. The system uses hidden Markov model [13] (Hidden Markov Model: HMM) to model words units and sentence of a language. In this work, the interest is to model the Moroccan dialect and implement a recognition system that converts a signal into a meaningful message that can be used thereafter. There is a several applications for a speech recognition system of Moroccan dialect. Most interesting are the man-machine dialogue, ie the passage of oral telephone calls; learning Moroccan dialect and systems helping people with disabilities [1]. The Moroccan dialect is a very important part of popular culture and covers almost the different regions.

The significance of ASR, several free systems have been developed, among the best known: HTK [11] and CMU Sphinx [2-3]. We used the last, it is based on Hidden Markov Model [3] and widely used in the field of speech recognition. In this context, our work focuses on the establishment of foundations for building a system of automatic speech recognition concerning the Moroccan dialect based on Sphinx4 [1].

In the following, we will outline the work done by starting with a theoretical approach to the hidden Markov model and dynamic programming (Section 2). Then, we present in brief a description of Moroccan dialect (section 5). The comparison results obtained from the hidden Markov model and dynamic

programming are given in Section 6. And it ends with a conclusion and outlook in Section 8.

II. THEORETICAL BASES

a) Hidden Markov Model

The hidden Markov model is a stochastic system capable [19], after a learning phase, to estimate the likelihood of observation sequence was generated by this model. The case represents a set of acoustic vectors of a speech signal. The hidden Markov model can be seen as a set of discrete states and transition between these states, it can be defined by all of the following parameters:

N : the number of model states

$A = \{a_{ij}\} = P(q_j/q_i)$: is a matrix of size $N * N$. It characterizes the transition matrix between states of the model. The transition probability to state j depends only on the state i :

$$P(q_t = j/q_{t-1} = i, q_{t-2} = k, \dots) = P(q_t = j/q_{t-1} = i) \quad (1)$$

$B = \{b_j(o_t)\} = P(o_t/q_j)$, where $j \in [1, N]$ is the set of emission probabilities of the observation o_t when the system is in the state q_j . The shape of this probability determines the type of HMM used. In this work, we use a continuous probability density [19] defined by the bellow relation:

$$b(o, m, v) = N(o, m, v) = \frac{1}{\sqrt{(2\pi)^n |c|}} e^{-\frac{1}{2}(o_n - m_i)c^{-1}(o_n - m_i)'} \quad (2)$$

Where:

O : Observation frame

C : covariance matrix (diagonal)

$$C = \frac{1}{n-1} \sum_{k=0}^n (o_k - m_k)' * (o_k - m_k)c$$

m : the mean of each coefficient

$$m = \frac{1}{n} \sum_{k=1}^n o_k$$

Taking into account several pronunciations of a word requires the use of a multi-Gaussian probability density [21] that the resulting probability is given by:

$$B_i(o_t) = \sum_{i=1}^k C_{ij} * b_j(o_t) \quad (3)$$

^{a Q β ψ ¥} : Traitement de l'information, Faculté des Sciences et Techniques PB 523, Béni Mellal, Maroc.

E-mails : hmadgm@yahoo.fr, daouic@yahoo.com,

najlae_idrissi@yahoo.fr, takfad@yahoo.fr, bbouikhlene@yahoo.fr.

k : number of Gaussian

C_{ij} : Gaussian weight of i in j

$B_j(o_t)$: observation probability at time t for state j

b) Dynamic programming

Dynamic programming [18] is one of the algorithms used in speech recognition domain, the principle is to compare two speech signals based on the distance between two matrices corresponding to the coefficients of Mel [18] of the two signals. Calculates the Euclidean distance between two vectors is sound by using the relationship:

$$d(i, j) = \sqrt{\sum_{k=1}^K (x_k - y_k)^2}, \quad i = \begin{pmatrix} X_1 \\ \vdots \\ X_N \end{pmatrix}, \quad j = \begin{pmatrix} Y_1 \\ \vdots \\ Y_N \end{pmatrix}$$

Then, calculate the minimum distance by traversing the element of the matrix obtained using the relation:

$$g(i, j) = \min \begin{cases} g(i-1, j) + d(i, j) \\ g(i-1, j-1) + 2 \cdot d(i, j) \\ g(i, j-1) + d(i, j) \end{cases}$$

The final distance is:

$$G = \frac{g(I, J)}{I+J}$$

Where:

I, J : Length acoustic arrays corresponding two signals.

III. EXTRACTION OF ACOUSTIC PARAMETERS

a) Pretreatment

The speech signals used were acquired using a microphone. The noise intra sentence was deleted manually using the tool wavsurfer. The digitized signals will be represented by a family $(x_n)_{n \in [1, k]}$ where k is the total number of samples. After, the signal is sampled using the computer's sound card with a frequency $F_s = 16\text{kHz}$ ie taking values follows a period $1/F_s$ seconds.

i. Mel coefficients

Parameterization of speech signals is to extract the coefficients of Mel. This stage is based on the Mel scale to model the perception of speech in a manner similar to the human ear, linear up to 1000 Hz and logarithmically above [22]. The importance of the logarithmic scale appears when using a broad bench of values as it helps to space the small value and approach large values. The digitized signals must be further processed for use in the recognition phase. To do this the pre-emphasis is performed to meet the high frequencies:

$$h_n = 1 - 0.97 * z_n^{-1} \quad (5)$$

Then the signal is segmented into frames each frame contains N sample of speech and includes almost 30ms of speech, to do this we use a sliding time window

of size 256. The successive windows overlap by half of their size ie 128 points in common between two successive windows. In this work we used the Hamming window [23]:

$$w(n) = 0.54 + 0.46 * \cos(2\pi * \frac{n}{N-1}) \quad (6)$$

In the next step the signal spectrum is calculated, it can introduce the signal (time domain) in frequency domain using the fast Fourier transform FFT:

$$X(n) = \frac{1}{N} \sum_{k=0}^{N-1} x(n) e^{jk 2\pi \frac{n}{N}} \quad (7)$$

To simulate the functioning of the human ear, we filter the signal through a bank of filters that each have a triangular response bandwidth. The filters are spaced so that their evolution is the Mel scale [22]. The approximate formula of the scale of Mel:

$$\text{Mel}(f) = 2595 * \log_{10}(1 + \frac{f}{700}) \quad (8-1)$$

$$X(i, k) = \sum_{n=0}^{N/2} X(n, k) * \text{Mel}(n, k) \quad (8-2)$$

The speech signal can be seen as the convolution in the time domain excitation signal $g(n)$ and the vocal tract impulse response $h(n)$:

$$x(n) = g(n) * h(n) \quad (9)$$

The application of the logarithm of the model on this equation gives:

$$\text{Log}|X(k)| = \text{Log}|G(k)| + \text{Log}|H(k)| \quad (10)$$

Finally, to obtain the coefficients of Mel applying the inverse Fourier transform defined by:

$$\text{FFT}^{-1}\{X(i, n)\} = x(n) = \frac{1}{N} \sum_{k=\frac{N}{2}}^{N-1} X(i, n) e^{jk 2\pi \frac{n}{N}} \quad (11)$$

This gives a vector of coefficients on each Hamming window. The number of filter adopted in this work is 12, it added the first and second derivatives of these coefficients, which gives in total 39 coefficients. Figure 1 gives a summary of the extraction of Mel coefficients (MFCC).

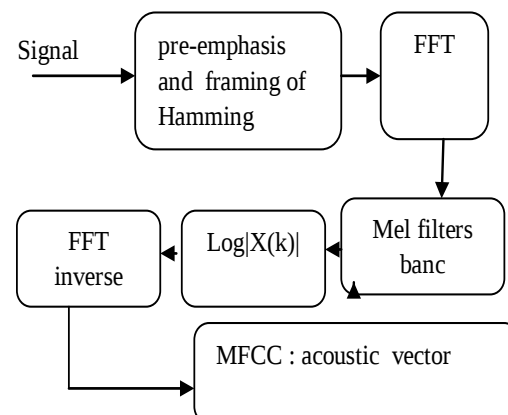


Fig 1 : Stage for acoustic parameters extraction

IV. LEARNING

After the extraction phase, the speech signal is represented by a matrix $N \times 39$ which N is the number of windows in the signal. The audio files used in the learning phase must be segmented into phonemes; each word corresponds to a sequence of phonemes. Each of these will be represented by a hidden Markov model with three states, each state is characterized by:

- Vector averages for a state i , is given by:

$$m_i = \frac{1}{n} \sum_{k=1}^n O_k, \quad n: \text{number of vectors for each state}$$

O_k : Observation vector number k .

- Covariance matrix for state i :

$$Coi = \frac{1}{n-1} \sum_{k=1}^n (O_k - m_i)' * (O_k - m_i) \quad (12)$$

The calculation of the mean vector and covariance matrix is performed for each Gaussian. In this paper we use five Gaussian so there will be five vehicles and five averages covariance matrices for each state. The calculation of the probability of resulting observation for each state is realized by the relationship 3.

Learning the model tends to maximize the logarithm of the probability of observation called the likelihood, to do this we use the Baum-Welch algorithm [15], whose steps are:

1- Initializing the model

- Creation of HMM for each state

- Initialization of the initial probability vector π with a higher probability for the first state and non-zero for the other two remaining states.

- Initialization of the transition matrix with probabilities respecting any transitions that the sum is equal to 1 and the model is a left-right (upper diagonal).

2 - Maximization: In this step, each iteration updates the model parameters and calculate again the likelihood. The updating of the model parameters is done via the following relations:

$$C_{ij} = \frac{\sum_{t=1}^T \gamma_t(j,k)}{\sum_{t=1}^T \gamma_t(j)}, \quad 1 \leq j \leq N, 1 \leq k \leq M \quad (13)$$

$$m_j = \frac{\sum_{t=1}^T o_t \gamma_t(j,k)}{\sum_{t=1}^T \gamma_t(j)}, \quad 1 \leq j \leq N, 1 \leq k \leq M \quad (14)$$

$$Co_j = \frac{\sum_{t=1}^T (o_t - m_j)(o_t - m_j)' \gamma_t(j,k)}{\sum_{t=1}^T \gamma_t(j)}, \quad 1 \leq j \leq N, 1 \leq k \leq M \quad (15)$$

With:

M : number of Gaussian.

N : number of acoustic vectors for each state.

With:

$$\gamma_t = \frac{\alpha_t(j)\beta_t(j)}{\sum_i \alpha_t(i)\beta_t(j)}$$

$$\gamma_t(j,k) = \gamma_t \left(\frac{C_{jk} N(o_t, m_{jk}, Co_{jk})}{\sum_{k=1}^M C_{jk} N(o_t, m_{jk}, Co_{jk})} \right)$$

C_{jk} is the weight of the Gaussian k relative to the state j and the coefficients α and β are calculated by the Forward-Backward algorithm [15].

V. RECOGNITION

The principle of recognition can be explained as the calculation of the probability $P(W/O)$: the probability that a sequence of words W is the signal S and to determine the word sequence that maximizes this probability.

According to Bayes formula the probability $P(W/S)$ can be written:

$$P(W/S) = P(w) \cdot P(S/W) / P(S) \quad (2)$$

With:

- $P(W)$: Prior probability of word sequence W : (Sample language).
- $P(S/W)$: Probability of signal S , given the sequence of words W (Acoustic Model).
- $P(S)$: probability of the acoustic signal S (independent of W).

The figure 2 shows the various stages of recognition, as a first step the signal is treated to extract acoustic vectors, based on these vectors the acoustic model is built from the HMM of phonemes learned on the training corpus. The succession of phonemes HMMS form the words models.

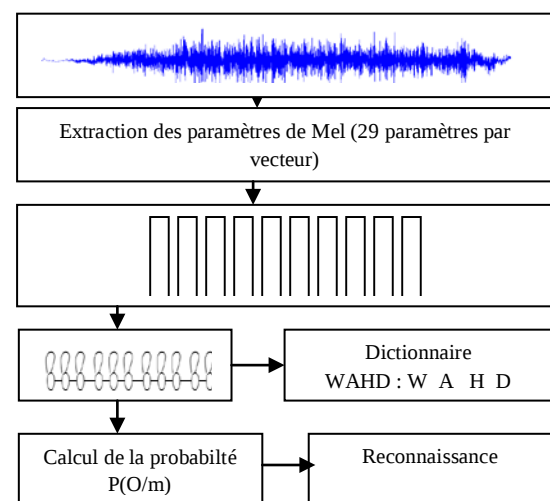


Fig. 2 : Stages of the recognition

VI. PRESENTATION OF MOROCCAN DIALECT

The Moroccan dialect called Darija is the popular language broadcast in almost all regions of the country. This dialect is a communication tool widely used and is different from one region to another. The dialect Darija contains almost Arabic words in addition to a regional component, the difference between classical Arabic and dialect Darija is at the pronunciation. The figure below shows an illustration of the difference:

bases	Pronunciation	Scripture	
The succession of two consonants (sokoun) is permitted. Two successive consonants come together.	OKTOB	اكتب	Classical Arabic
Most letters are pronounced with 'sokoun'. Most words are pronounced without vowels (SAKINA)	KTB	كتب	Darija

Fig 1. Difference between classical Arabic and Darija

VII. EXPERIMENTAL RESULTS

a) Learning base

The learning base used in our system contains 2500 pronunciation, the characteristics of the training set are illustrated in the following table:

Duration of the training set	Number of pronunciations
1h40min of pronunciations.	2500 recorded pronunciation independently and in different situations

Table1 : Characteristics of the learning base

The construction of the training set was made by taking the pronunciation of Arabic numerals 0 to 9 in the Moroccan dialect, Table 2 shows the formation of the learning base.

numbers	Phonetic Transcription
0	S I F R
1	W A H D
2	J U J
3	T L A T A
4	R B 3 A
5	X M S A
6	S T T A
7	S B 3 A
8	T M N Y A
9	T S 3 U D

Tab.2 : Phonetic transcription used for the recognition of digits in dialect Darija

b) Results

The test database contains 300 different pronunciations including noisy audio files. The recognition quality is measured by calculating the rate of recognition given by equation (3):

$$t = \frac{\text{number of words recognized}}{\text{size of the test database}}$$

The results obtained are shown in Table 4.

Test database	Results
300 different pronunciations introducing more noisy audio files	T=91%

Tab.4 : Results for the recognition system of the dialect Darija

The comparison of the results was made on noisy audio data. Table 5 illustrates the results obtained.

	HMM	DTW
Execution Time	Very fast	Plus then 10s for a big wav files
Recognition rate	91%	60%

Tab.5 : Results of comparison between the HMM and DTW

The efficiency of dynamic programming appears on the audio files not noisy. The disadvantage is that the execution time increases proportionally with the length of the file, which influence the time of recognition. In comparison with dynamic programming, hidden Markov model can model a word by a sequence of phonemes and sentence by a sequence of word models, which makes this process more effective and more appropriate to be implemented in systems Recognition advanced.

VIII. CONCLUSION

This work enables the establishment of a voice recognition system of the Moroccan dialect. This article can give an idea about the phonetics used for the recognition of the language. In comparison with dynamic programming, the results obtained by the

hidden Markov model are very satisfactory despite the limited number of speakers and size of the database. This shows the importance of stochastic and probabilistic modeling in the field of recognition.

Based on what has been achieved in this work, we'll build a system of passing oral phone call on the Moroccan dialect integrated into mobile phones, helping people with disabilities and people who do not dial telephone numbers.

REFERENCES REFERENCES REFERENCIAS

1. Ali sadiqui & Noureddine chenfour "Reconnaissance de la parole arabe basé sur CMU Sphinx", Séria Informatica. Vol VIII fasc. 1 2010.
2. H. Satori & M. Harti "Système de la reconnaissance de la reconnaissance automatique de la parole", Faculté des Sciences, B.P. 1796, Dhar Mehras Fès, Maroc.
3. Cornijeol and L. Miclet, "Apprentissage artificielle-méthode et concept" 1988.
4. T. Pellegrini et Raphael, "Durée suivi de la voix parlée garce au modèle caché" 1989.
5. R. Gonzales and M. Thomson, "Syntactic pattern recognition" 1986.
6. Divejver and J. Killer, "Pattern recognition" in Pattern Recognition: a statistical approach"; Prentice Hall 1982.
7. Reweis, "Hidden Markov-Modele-Sam" 1980.
8. Robiner and Juang. "Fundamentales of speech recognition" 1993.
9. M. Amour ,A. Bouhjar & F. Boukhris IRCAM: publication : "initiation à la langue Amazigh" 2004.
10. RAP: Thèse 2008-[Benjamin LECOUEUX].
11. B. Resch "Automatic Speech Recognition with HTK" 2003.
12. Chunsheng Fang (2009) "From Dynamic Time Warping (DTW) to Hidden Markov Model" (HMM) University of Cincinnati
13. A. Cornijeol and L. Miclet, "Apprentissage artificielle-méthode et concept" 1988.
14. P. Galley, B. Grand & S. Rossier, "reconnaissance vocale Sphinx-4" EIA de Fribourg mai 2006.
15. T. Pellegrini et R. Duée "Suivi de la voix parlée grâce aux modèles de Markov Caché", lieu : IRCAM 1 place Igor Stravinsky 75004 PARIS juin 2003.
16. S. Sigurdsson, Kaare Brandt Petersen and Tue Lehn-Schiøler "Mel Frequency Cepstral Coefficients: An Evaluation of Robustness of MP3Encoded Music", Informatics and Mathematical Modelling Technical University of Denmark Richard Petersens Plads - Building 321 DK-2800 Kgs. Lyngby - Denmark.
17. T. AL ANI "Modèles de Markov Cachés (Hidden Markov Models (HMMs))", Laboratoire
18. A2SI-ESIEE-Paris / LIRIS.
19. G. SEMET & G. TREFFOT "La reconnaissance de la parole avec les MFCC" TIPE juin 2002.
20. A. Chan, Evandro Gouvêa & Rita Singh "Building Speech Applications Using Sphinx and Related Resources": <http://docpp.sourceforge.net> , August 2005.
21. Dr. A. Drygajlo "Introduction aux statistiques gaussiennes et à la reconnaissance statistique de formes", Ecole Polytechnique Fédérale de Lausanne.
22. S. Jamoussi, "Méthodes statistiques pour la compréhension automatique de la parole", Ecole doctorale IAEM Lorraine, 2004.
23. SEMET Gaetan & TREFFO, 'Grégory,'Reconnaissance de la parole avec les coefficients MFCC' TIPE juin 2002.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Fine Grained Access Control Model Based on Diverse Attributes

By Rajender Nath , Gulshan Ahuja

Kurukshetra University, Kurukshetra, Haryana, India

Abstract - As the web has become a place for sharing of information and resources across varied domains, there is a need for providing authorization services in addition to authentication services provided by public key infrastructure (PKI). In distributed systems the use of attribute certificates (AC) has been explored as a solution for implementation of authorization services and their use is gaining popularity. AC issued by attribute authority (AA) facilitates identification of a service requester and can be used to enforce access control for resources. AC of a service requester is used as part of credentials supplied during the service request for accessing any resource. As there exist potentially multiple issuing domains which issue credentials, therefore the target domain must allow access to resources by considering different credentials and must be able to decide about which set of attributes can be considered as valid attributes for making access control decisions. In this paper, we present an authorization based access control model that allows a fine grained access control to resources in an open domain by utilizing attributes issued by diverse attribute authorities.

Keywords : *attribute certificates, attribute authority, authorization, access control.*

GJCST Classification : *D.4.6*



Strictly as per the compliance and regulations of:



A Fine Grained Access Control Model Based on Diverse Attributes

Rajender Nath^α, Gulshan Ahuja^Ω

Abstract - As the web has become a place for sharing of information and resources across varied domains, there is a need for providing authorization services in addition to authentication services provided by public key infrastructure (PKI). In distributed systems the use of attribute certificates (AC) has been explored as a solution for implementation of authorization services and their use is gaining popularity. AC issued by attribute authority (AA) facilitates identification of a service requester and can be used to enforce access control for resources. AC of a service requester is used as part of credentials supplied during the service request for accessing any resource. As there exist potentially multiple issuing domains which issue credentials, therefore the target domain must allow access to resources by considering different credentials and must be able to decide about which set of attributes can be considered as valid attributes for making access control decisions. In this paper, we present an authorization based access control model that allows a fine grained access control to resources in an open domain by utilizing attributes issued by diverse attribute authorities.

Keywords : attribute certificates, attribute authority, authorization, access control.

1. INTRODUCTION

With continually changing business environment privacy and protection of resources is becoming more and more important. The access control to resources is bound up with the authentication and the authorization. There is a strong need felt for receptive authorization infrastructure that can cater for rapidly changing dynamic environments and should be able to validate the identity of service requesters.

The commonly used credentials for access are identity credentials, attribute credentials and authorization credentials. Identity based access control systems require identity certificates which are issued and certified by certification authorities (CA). When a CA issues an identity certificate, it binds a particular public key to the name of the service requester (SR) identified by the certificate. In addition to a public key, a certificate always includes information such as the validity period, the name of the CA, the digital signature of the issuing CA etc. Identity based credentials are more suitable

where service requesters are already known to the service provider (SP) through the process of registration. This approach works well in a tightly coupled environment. Identity based access control puts a constraint of prior registration of every service requester which limits the scalability of overall system. Another access control approach is based on the authorization certificates which are based on the principle of delegation of rights and responsibilities. The authorization certificates are issued by the authorization authorities who have rights to access the specific resource and thus can delegate full or subset of rights to other users. The authorization certificates usually contain the identity of the resource, identity of the service requester, access rights to access the resource, etc. The advantage of authorization certificates are that service requesters are authenticated in their own domain and another service requester to whom the rights have been delegated can realize the access control based on delegated rights. A different area of research developments is access control based on attributes. In attributes based access control systems the access policy is based on the various attributes which are assigned to the service requesters. The attribute certificates are issued by Attribute Authority (AA) and these contain the name value pairs of the various attributes. Attributes based authorization offers more flexibility and scalability for an open and distributed environment. The use of AC based on privilege management infrastructure (PMI), allows including and revoking attributes and can contain information about the privileges or roles of a user. AC conveys a short-lived attribute about a given subject and can be used to authenticate the identity of the attribute certificate holder. A real time problem arises when the request made by a service requester requires attributes which have been issued by diverse attribute authorities and are located at different locations. The authorization efforts become more difficult when two or more AAs save attributes for a service requester with different identities. The rest of this paper is structured as follows. Section II highlights the related work. Section III highlights the requirements to develop a new model based on diverse attributes. Section IV describes the implementation architecture and explains the working of proposed model. Finally Section V concludes and briefly describes scope for the future work.

Author ^α : Department of Computer Science & Applications,
Kurukshetra University, Kurukshetra, Haryana, India.
E-mail : rnath_2k3@rediffmail.com

Author ^Ω : Tilak Raj Chadha Institute of Management & Technology,
Yamunanagar, Haryana, India.
E-mail : ahujag_24@yahoo.com

II. RELATED WORK

Traditional access control approaches base their authorization decisions on subject's identity. A number of research papers based on attributes based authorization have been proposed by researchers. Ioannis Mavridis et al. [1] proposed a mechanism for access control based on attribute certificates for medical Internet applications. David Chadwick [2] proposed X.509 privilege management infrastructure. Later David Chadwick et al. [3] proposed Role-Based Access Control with X.509 Attribute Certificates. The proposed approach in paper adopted the standard X.509 PMI to build an efficient role-based trust management system in which role assignments can be widely distributed among organizations, and an XML-based local policy determines which roles to trust and which privileges to grant. Jordi Forne et al. [4] presented an implementation of an authorization system for web based applications based on the ITU-T X509 recommendations which specifies use of privilege management infrastructure for realizing access control. Access control mechanism based on authorization, using attributes issued by a remote attribute authority, has been proposed by S. Cantor. [5]. Wei Zhou et al. [6] proposed a role based access control with attribute certificates. Alfieri R. et al. presents a VOMS model [7] for managing authorization in a Grid Environment and allows coalition of multiple attributes. Eric Yuan [8] proposed an attribute based access control (ABAC) model as a new approach, which is based on subject, object, and environment attributes and supports both mandatory and discretionary access control needs. M Liu et al. [9] proposed an attribute and role based access control model ARBAC for web services. However, the role remains static and when assigned it becomes out of date. Alan H. Karp [10] proposed an implementation based on authorization based access control (ABAC) for services oriented architecture. David W Chadwick [11] presented a model and protocol elements for linking AAs, service providers and user attributes together, under the sole control of the user and allowed merging the attributes from multiple AAs in order to grant the user access to its resources. Frikken K et al. [12] proposed an approach for attribute based access control with hidden policies and hidden credentials. Shen Hai Bo et al. [13] proposed an attribute based access control model for web services. Nirmal Dagdee et al. [14] proposed an access control methodology for sharing of open and Domain confined data using Standard Credentials. The methodology requires that various types of standard credentials and related attributes are identified and published by some apex authority so that the resource providers can define their access policies in terms of these standard credentials. In real terms, identification of standard credentials is a very difficult task and is not suitable for

largely distributed systems having millions of service requesters. Regina N. Hebig et al. [15] describe a prototype implementation with an architecture based on the standards XACML, SAML, WSPolicy, WS-SecurityPolicy and WS-Trust which puts the focus on sharing identity and attribute information across independent domains for the purpose of access control.

III. PROBLEM FORMULATION

The service requester's credentials may be stored or issued in a variety of places, for example, each AA may store the attributes or the credentials it issues in its own repository. When a service requester makes a request to service provider for accessing a resource and presents its credentials. At service provider's end, the presented set of attributes may not be sufficient enough to grant access to a resource. This necessitates that service requester must collect together the credentials required for making access to a resource. David W Chadwick [11] presented a model and protocol elements for linking attributes from multiple AAs. His approach requires input from the user who wish to link attributes from multiple AAs. However, the main issue with his approach is that user has to initiate multiple browser instances and execute steps for cross linkages between multiple attribute authorities. If the number of attributes required for grant of access belongs to multiple attribute authorities, the same process is to be carried multiple times for providing linkages between multiple attribute authorities. This makes the task of service requester more complex and time consuming. We propose a new model where the service requester's task of creating linkages is eliminated and the process of linkage is initiated by service provider only. The proposed work also takes care of the privacy concern of the service requester to ensure that service provider will be able to link attributes from multiple attribute authorities only when service requester desires to create linkages with multiple attribute authorities.

IV. PROPOSED MODEL

This section describes the details of proposed model. The approach requires that all organizations who are willing to exchange and share information among them must form agreements for a number of conditions i.e. security mechanisms to be used, attribute definitions etc. and must pre establish a certain level of trust. Such sort of arrangement is termed as federation and is same as Shibboleth federations [16]. The mechanism assumes that linking between AAs is based on secured shared information and there also has to be secured shared information between a service requester and all attribute authorities where service requester is registered for attribute sharing.

An authorization model based on diverse attributes from different attribute authorities is shown in

Figure 1. The diagram reflects following components involved in the access mechanism.

- Policy Enforcement Point (PEP) intercepts all incoming requests from service requesters. Every received request is sent to the PDP. PEP allows or denies access to the service requester for access of requested resource.
- Policy Decision Point (PDP) evaluates the applicable policies against service requests. PDP checks for the available attributes in the service request to check whether access request can be granted or not. In case the attributes contained in service request are not sufficient enough for grant of access to resource, it hands over the request parameters along with information about additional attributes required for grant of access request.
- Attributes authorization and linker module (AALM) is responsible for contacting concerned attribute authorities for making request of additional

attributes and sends back the attributes received from multiple attribute authorities to PDP.

- Policy store contains policies for making access control decisions. The policies can be stored in XML format as it allows standard representation of access control rules. Extensible access control markup language (XACML) [17] [17] can be used to allow implementation of access control policies.
- Policy management interface allows handling of policies in the policy store.

a) Overview

A service requester can acquire multiple identities by registering with a number of attribute authorities. The decision to register with the attribute authority can be based on its reputation, quality of service etc. The mechanism requires that every service requester and all AAs to whom requester wishes to link must exchange and agree upon conversation framework for transfer of information between them.

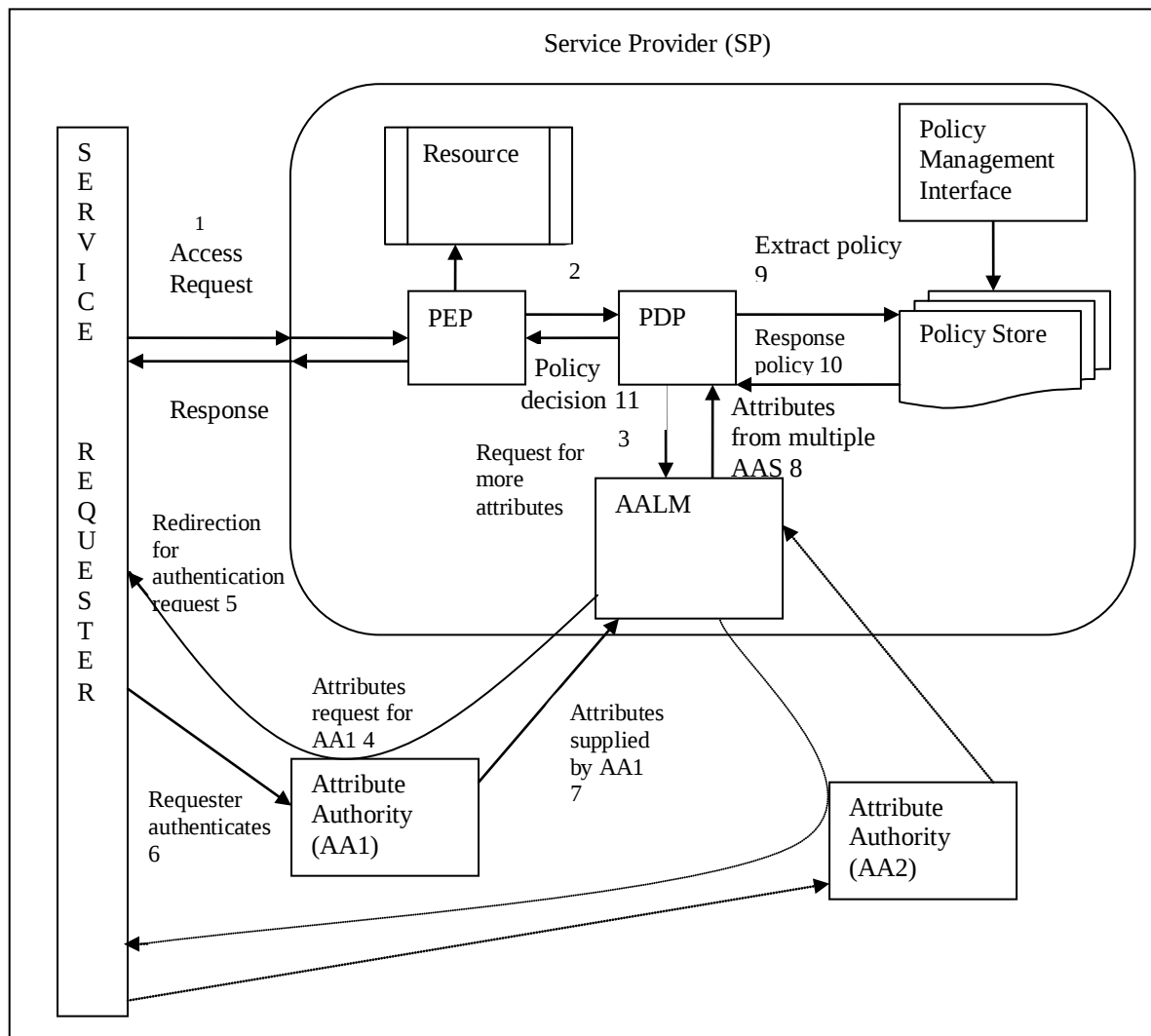


Figure 1 : Diverse attributes based authorization model

The service provider should not be able to obtain attributes from any AA without knowledge and permission from the service requester. The AAs who are willing to exchange information must make groups with predefined policies and rules. There has to be one or more than one primary attribute authority. The primary attribute authorities act as the root for all other AAs which are members of the group. Before making a request for accessing any resource, the service requester must acquire credentials from one of the primary certified authority. The attributes returned by the primary AA contain a basic set of the attributes along with information about all AAs for which service requester has already registered and has agreed to use additional attributes. Each service provider in the federation is free to decide about the number and types of attributes for granting access requests. The service providers may grant access on basic set of attributes or may decide to impose more security check by imposing requirement for additional attributes from one or more AAs.

b) SAML based conversation framework

We use SAML assertions for describing conversation tokens. Figure 2 depicts conversation framework between service requester and provider. Figure 3 describes the format of SAML based conversation tokens.

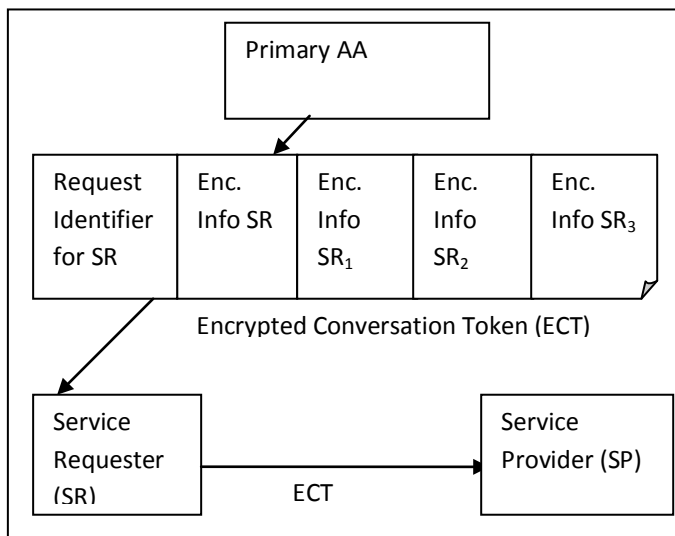


Figure 2 : Conversation Framework

Let ATS_b be the basic set of attributes and Σ is an alphabet, a non-empty finite set.

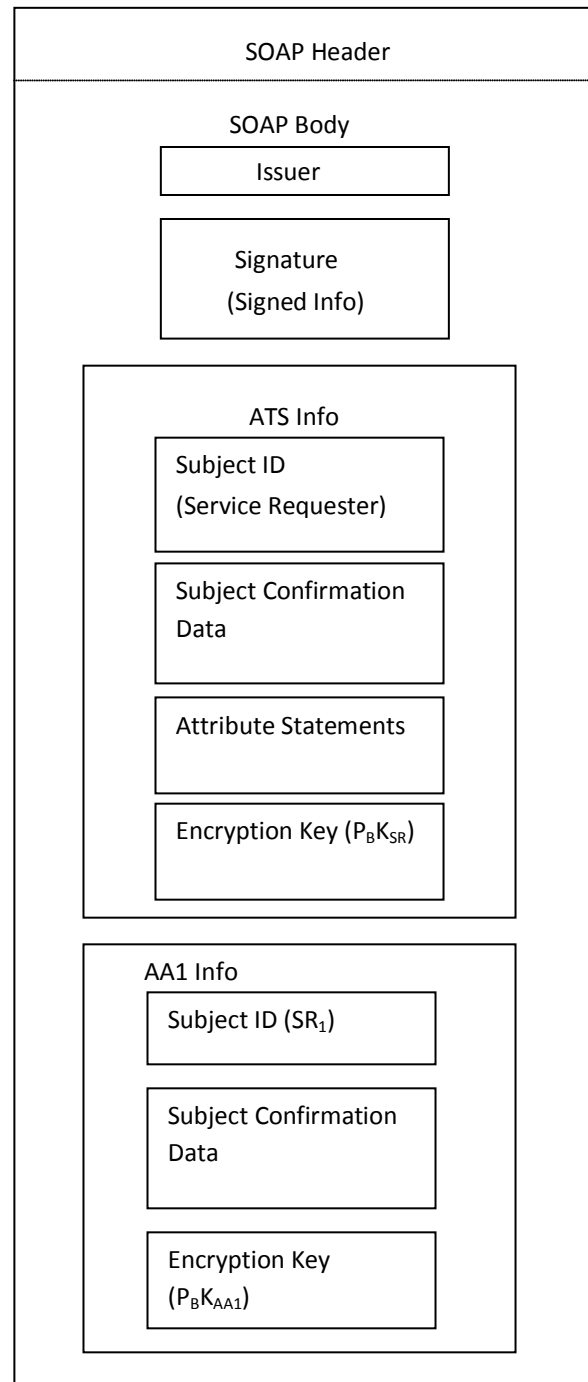


Figure 3 : Conversation Token Format

The set of all attributes over Σ of any length can be denoted in terms of Σ^n as

$$ATS_b = \Sigma^* = \bigcup_{n \in \mathbb{N}} \Sigma^n$$

Let service requester SR is registered with three attribute authorities. As identity SR_1 with attribute authority AA1, as identity SR_2 with AA2, as identity SR_3 with AA3. The conversation token from the primary attribute authority to service requester will be of the form as defined below.

$ATS\{SR, Time_Stamp\}P_{BK_{SR}}$
 $AA1\{SR_1, Time_Stamp\}P_{BK_{AA1}}$
 $AA2\{SR_2, Time_Stamp\}P_{BK_{AA2}}$
 $AA3\{SR_3, Time_Stamp\}P_{BK_{AA3}}$

c) Authentication of conversation tokens

As per figure 1 and figure 2, the authorization process is divided in to 2 parts: obtaining of basic set of attributes from the primary attribute authority and use of attributes for making authorization decisions. For the first part, since it is assumed that service requester trusts the primary attribute authority and can decrypt and obtain the basic set of attributes using its private key $P_{VK_{SR}}$. The primary attribute authority also passes on the encrypted info for every other AA where service requester's attributes are already located. This encryption is carried using public key of the corresponding attribute authority. The second part is an establishment stage where ECT is used by service provider to make access decision. When a service requester makes a request for accessing a resource, the service provider executes following tasks:

Task 1: Service requester obtains credentials from primary attribute authority. The credentials issued by primary attribute authority to the service requester are encrypted using public key of respective authority and are sent in format as in figure 3.

Task 2 :The service requester decrypts the basic attributes, using its private key $P_{VK_{SR}}$ and presents the basic set of attributes along with encrypted info about AAs, acquired from primary AA, to service provider.

Task 3 :On receiving the request, the service provider checks for the basic set of attributes against already specified policies in the policy store to decide whether access can be granted or not.

Task 4 :In case the existing policies do not allow access based on basic set of attributes contained in the service request, the PDP module passes the service request to the AALM module.

Task 5 :AALM module extracts the information from the service request to find out for which all other AAs service requester has already registered. The request message along with the requester's URL was encrypted using public key of concerned AA so it can be decrypted using private key by the concerned AA only. The service provider just knows that at which attribute authority the service requester is registered so this helps to maintain the privacy concern of the requester because service provider can not determine the identity and attributes of the requester located on a particular AA.

Task 6 : AALM sends a request message for the required attributes to the concerned AAs along with the encrypted info about requester's identity and date time stamp.

Task 7 : AA decrypts the information corresponding to its requested attribute set using its private key and extracts the identity of service requester and also verifies the date and time stamp for validity of message. For example SR had already registered with AA1 as SR_1 , therefore upon successful decryption of the message, AA1 can ascertain about SR_1 .

Task 8 : To make sure that SR is willing to allow attributes from AA, it redirects an authentication request to SR.

Task 9 : Once the service requester authenticates with AA, the information regarding attributes required by the service provider is shown and service requester is given a choice to allow passing back the set of attributes to the service provider.

Task 10 : Once the confirmation is made by the service requester, the one or more required attributes are sent back to the service provider.

Tasks from 6 to 9 are repeated for every AA to whom AALM module sends a request for additional attributes. In the event of any AA failing to provide required attributes the request is terminated with an appropriate response to the service requester.

The access control decision based on diverse attributes can be realized in terms of a function

$$f(ATs_b) \text{ or } f(ATs_b \times ATs_i \times ATs_j \times ATs_k)$$

Where ATs_b is the basic set of attributes for SR, and ATs_i , ATs_j , ATs_k are the set of attributes for three different attribute authorities identified as AA_i , AA_j , AA_k respectively.

The above mentioned function is implemented and used by PEP component to decide whether the access to resource can be allowed based on basic set of attributes or attribute assignments from multiple AAs can be evaluated. The implementation of function solely depends upon the policies and requirements of the service provider. The evaluation outcome can be considered for granting access to the resource. The access control mechanism discussed in this paper allows in implementing fine grained access control based on multiple attributes. The proposed approach allows every service provider to decide the level of security for granting access request. The service provider can choose to allow access request based on the basic set of attributes or may put more restrictions by imposing requirements for attributes from one or more AAs. The use of basic set of attributes works well for a large number of service requesters who most of the times need access to resource based on basic information. However, the ability for any service provider to impose requirement for additional attributes helps in achieving fine grained access control.

V. CONCLUSION

In this paper, authors have proposed a mechanism for allowing access to a resource based on the multiple attributes from one or more AAs. The merit of the proposed approach is that service provider can link to the attribute authorities and obtain attributes for grant of access only when it is permitted by the service requester. Even if the service provider is able to determine that to which all AAs the service requester has already registered, it can not automatically obtain attributes without service requester's permission. The proposed approach focuses only on authorization of requests based on attributes. The future work may consider other aspects related with attributes based access without involvement of centralized authority and automated trust establishment.

REFERENCES REFERENCES REFERENCIAS

1. Ioannis Mavridis, Christos Georgiadis, George Pangalos, marie Khair, "Access Control based on Attribute certificates for Medical Internet applications", Journal of medical Internet Research, Vol 3, 2001.
2. David Chadwick, "The X.509 Privilege Management Infrastructure", <http://sec.cs.kent.ac.uk/download/X509pmiNATO.pdf>, 2002 .
3. David W. Chadwick, Alexander Otenko, and Edward Ball, "Role-Based Access Control with X.509 Attribute Certificates", IEEE internet computing, march-april 2003, pp. 62 – 69.
4. J. F d an d M. F. Hanarejos , "Web-based Authorization based on X.509 Privilege Management Infrastructure ",IEEE Pacific Rim Conference on Communications, Computers and signal Processing, 2003.
5. S. Cantor. "Shibboleth Architecture, Protocols and Profiles", Working Draft 02. 22 September 2004, <http://shibboleth.internet2.edu/>
6. Wei Zhou, Christoph Meinel, "Implement Role based access control with attribute certificates", The 6th IEEE International Conference on Advanced Communication Technology, page(s): 536- 540, 2004.
7. Alfieri, R., Cecchini, R., Ciaschini, V., Dell'Agnello, L., Frohner, A., Lorente, K., Spataro, F., "From gridmap-file to VOMS: managing authorization in a Grid environment". Future Generation Computer Systems. Vol. 21, no. 4, pp. 549-558. Apr. 2005.
8. Eric Yuan et al., "Attributed Based Access Control (ABAC) for Web Services", IEEE International Conference on Web Services 2005.
9. Liu M., GUO H Q., and SU J D., "An Attribute and Role-Based Access Control Model for Web Services", In proceedings of The Fourth International Conference on Machine Learning and Cybernetics, Guangzhou, China, 18-21 August 2005.
10. Alan H. Karp, "Authorization-Based Access Control for the Services Oriented Architecture, Proceedings of the Fourth International Conference on Creating, Connecting and Collaborating through Computing", 2006.
11. David W Chadwick, "Authorisation using Attributes from Multiple Authorities", Proceedings of the 15th IEEE International Workshops on Enabling Technologies Infrastructure for Collaborative Enterprises 2006.
12. Frikken K, Atallah M, Jiangtao Li, "Attribute-Based Access Control with Hidden Policies and Hidden Credentials", IEEE Transactions on Computers, Volume 55, Issue 10, Page(s): 1259 – 1270, Oct. 2006.
13. Shen Hai Bo, Hong Fan, "An attribute based access control model for web services", Proceeding of the 7th International Conference on Parallel and Distributed Computing, Applications and Technologies, IEEE 2006.
14. Nirmal Dagdee, Ruchi Vijaywargiya, "Access control methodology for sharing of open and Domain confined data using Standard Credentials", International Journal on Computer Science and Engineering Vol.1(3), 2009, 148-155.
15. Regina N. Hebig et al., "A Web Service Architecture for Decentralized Identity- and Attribute-based Access Control", IEEE International Conference on Web Services, 2009.
16. <http://www.educause.edu/EDUCAUSEQuarterly/EDUCAUSEQuarterlyMagazineVolum/FederatedSecurityTheShibboleth/157315>
17. <http://www.oasis-open.org/committees/xacml/repository/>



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Recognition of Handwritten Tifinagh Characters Using a Multilayer Neural Networks and Hidden Markov Model

By B. El Kessab, C. Daoui, K. Moro, B. Bouikhalene , M. Fakir

Faculty of science and technology PB 523, Beni Mellal, Morocco

Abstract - In this paper, we propose a system for recognition handwritten characters Tifinagh, with the use of neural networks (the multi layer perceptron MLP), the hidden Markov model (HMM), the hybrid Model MLP/HMM and a feature extraction method based on mathematical morphology, this method is tested on the database of handwritten isolated characters Tifinagh size consistent (1800 images in learning and 5400 test examples). The recognition rate found is 92.33%. The MLP, HMM and MLP+HMM classifiers show good enough results.

Keywords : *Recognition of handwritten characters, Tifinagh characters, Neural Network, Hidden Markov Model, Hybrid Model, mathematical morphology.*

GJCST Classification : *G.3, F.1.1*



Strictly as per the compliance and regulations of:



Recognition of Handwritten Tifinagh Characters Using a Multilayer Neural Networks and Hidden Markov Model

B. El Kessab^α, C. Daoui^α, K. Moro^β, B. Bouikhalene^ψ, M. Fakir[¥]

Abstract - In this paper, we propose a system for recognition handwritten characters Tifinagh, with the use of neural networks (the multi layer perceptron MLP), the hidden Markov model (HMM), the hybrid Model MLP/HMM and a feature extraction method based on mathematical morphology, this method is tested on the database of handwritten isolated characters Tifinagh size consistent (1800 images in learning and 5400 test examples). The recognition rate found is 92.33%. The MLP, HMM and MLP+HMM classifiers show good enough results.

Keywords : Recognition of handwritten characters, Tifinagh characters, Neural Network, Hidden Markov Model, Hybrid Model, mathematical morphology.

I. INTRODUCTION

Automatic recognition of handwritten characters is the subject of several research for several years. This has several applications: In the field of multimedia and the compression of image etc... Automatic recognition of a Tifinagh character is done in three steps: the first phase is that of pretreatment to reduce noise, the second for extraction of characteristics and the third to make the classification (Neural Networks, Hidden Markov Model, Hybrid Model MLP/HMM). Neural networks are a system of calculate widely used for the recognition of images [1, 2, 3, 4]. In learning we use the gradient descent algorithm, in Hidden Markov Model we use a vector of extraction as a suite of observation and we seeks to maximize the model with the best probability. The Baum-Welch algorithm is used to learning. And for the Hybrid model MLP/HMM we considered the output of the neural networks as a probability of emission for the hidden markov model. We propose a system of recognition implemented on a base of Tifinagh characters manuscripts. This paper is organized as follows: Section 2 is devoted to the test database. In section 3, we describe the method of characteristics extraction based on mathematical morphology. In section 4, we present an overview on the neural networks MLP. Experimental

results of neural networks are presented in section 5. The hidden Markov model HMM is presented in section 6. Experimental results on the hidden Markov model will be presented in section 7. In section 8, we present the hybrid model MLP/HMM. Recognition system is illustrated in the following figure (Figure 1).

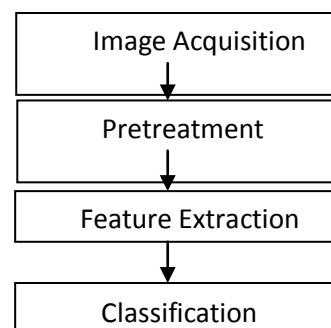


Figure 1: The recognition process

II. TIFINAGH DATABASE

The database used is Tifinagh, it is composed of 7200 characters Tifinagh (manuscripts + Imprimed) (1800 in learning and 5400 in test). The number of characters according to IRCAM Tifinagh is 33 characters. Example of IRCAM Tifinagh characters:



Figure 2 : Tifinagh characters

III. P REPROCESSING

In the preprocessing phase, the images of handwritten characters are rendered digital with a scanner, then it make the binary image with thresholding. After standardization for image extracted in a square the size standard is: 150 x 150.

Author ^α : Laboratory of modelling and calculation - Faculty of science and technology PB 523, Beni Mellal, Morocco,
E-mails : bade10@hotmail.fr, daouic@yahoo.com

Author ^{β ψ ¥} : Team Information Processing, Faculty of Science and Technology, PB 523, Beni Mellal, Morocco,
E-mails : kamalmoro@hotmail.com, bbouikhalene@yahoo.fr, fakfad@yahoo.fr

IV. EXTRACTION

The method used for extraction is based on mathematical morphology [5, 6, 7]. We seek to detect five zones characteristics for each image: West zone, East zone, North zone, South zone and central zone. These characteristic areas are detected by the dilatation of the processed image in the four directions.

a) The dilation of the image:

The dilation is a transformation based on the intersection between the object of the image A (white pixels) with a structuring element B. It is defined by the following formule

$$\text{Dilatation } (A, B) = \{x \in \text{Image} / B_x \cap A \neq \emptyset\}$$

Where A is the object of the image (the white pixels), B the structuring element which is a particular set of Center x, known size and geometry (in this work is a right half).

Example of the dilatation characters to the East.



Figure 3 : The dilatation characters to the East

And the same thing for other directions West, North and South.



Figure 4 : The dilations of the character to the West, North and South

b) Detection the characteristic zones of the image:

Is determined for each character the discriminating parameters (zones). The characteristic zones can be detected by the intersections of dilations found to the East, West, North and South. We define for each image five types of characteristic zones: East, West, North, South, and Central zone.

i. Extraction of East characteristic zone :

A point of the image (Figure 5) belongs to the East characteristic zone (Figure 6) if and only if:

- This point does not belong to the object (the white pixels in image).
- From this point, moving in a straight line to the East, we do not cross the object.
- From this point, moving in a straight line to the south, north and west one crosses the object (Figure.5). The result of the extraction is illustrated in (Figure.6).



Figure 5 : Image of the character



Figure 6 : The East characteristic zone (EZ)

ii. Extraction of West characteristic zone:

A point of the image (Figure 7) belongs to the West characteristic zone (Figure 8) if and only if:

- This point does not belong to the object (the white pixels in image).
- From this point, moving in a straight line to the West, we do not cross the object.
- From this point, moving in a straight line to the south, north and East one crosses the object (Figure.7). The result of the extraction is illustrated in (Figure.8).



Figure 7: Image of the character



Figure 8 : The West characteristic zone (WZ)

iii. Extraction of South characteristic zone:

A point of the image (Figure 9) belongs to the South characteristic zone (Figure 10) if and only if:

- This point does not belong to the object (the white pixels in image).
- From this point, moving in a straight line to the South, we do not cross the object.
- From this point, moving in a straight line to the north, East and West one crosses the object (Figure.9). The result of the extraction is illustrated in (Figure.10).



Figure 9 : Image of the character



Figure 10 : The South characteristic zone (SZ)

iv. Extraction of North characteristic zone :

A point of the image (Figure 11) belongs to the North characteristic zone (Figure 12) if and only if:

- This point does not belong to the object (the white pixels in image).
- From this point, moving in a straight line to the



Figure 11 : Image of the character



Figure 12 : The North characteristic zone (NZ)

v. *Extraction of Central characteristic zone :*

A point of the image (Figure 13) belongs to the Central characteristic zone if and only if:

- This point does not belong to the limit of the object.
- From this point, moving in a straight line to the south, north, east and west we cross the object. The result of the extraction is illustrated in the (FIG.13).



Figure 13 : Image of character



Figure 14 : The central characteristic zone (CZ)

Each character is characterized by five components: NWZ, NEZ, NNZ, NSZ and NCZ. with these latter parameters are the numbers of pixels of value 1 respectively in the characteristic zones West, East, North, South and Central. The vector of extraction will be defined as follows:

$$\text{Vect} = [\text{WZ}, \text{EZ}, \text{NZ}, \text{SZ}, \text{CZ}]$$

With Npixels is a Number of pixels in the image size 150 x 150

$$\begin{aligned} \text{WZ} &= \text{NWZ} / (\text{Npixels}). \\ \text{EZ} &= \text{NEZ} / (\text{Npixels}). \\ \text{NZ} &= \text{NNZ} / (\text{Npixels}). \\ \text{SZ} &= \text{NSZ} / (\text{Npixels}). \\ \text{CZ} &= \text{NCZ} / (\text{Npixels}). \end{aligned}$$

V. NEURAL NETWORKS

The neural networks [8, 9, 10, 11, 12] based on properties of the brain to build systems of calculation best able to resolve the type of problems as human beings live know resolve. They have several models, one of these models is the perceptron.

a) *Multi-layer perceptron (MLP)*

[13, 14, 15] The classification phase is as follows:

The number of neurons in the network is:

- Five neurons in the input layer (the number five

corresponds to the values found in the vector of extraction).

- Eighteen neurons in the output layer (the number eighteen corresponds to the number of characters used).
- The number of neurons in the hidden layer is selected according to these three conditions:
- Equal number of neurons in the input layer.
- Equal to 75% of number of neurons in the input layer. Equal to the square root of the product of two layers of exit and entry.

According to these three conditions are varied the number of neurons of layer hidden between five and ten neurons. The method used for learning is the descent of the gradient (gradient back-propagation algorithm) [16, 17, 18, 19, 20].

Layers	Neurons	Constant of learning
Input	5	$\alpha = 0.9$
Hidden	9	
Output	18	
Squared error		Activation function
$E = \frac{1}{2} (t - o)^2$ t : the theoretical output. o : the desired output.		Sigmoid function $F(x) = \frac{1}{1 + e^{-x}}$

Figure 15 : Details of the neural networks

VI. CLASSIFICATION

After extracting the characteristic data of the input image, we will classify its data with neural networks (Multilayer perceptron). We must compute the coefficients of weight and desired outputs.

a) *Experimental Results*

The values of characteristic vectors obtained in the extraction phase are introduced at the input of the neural network, and we know the desired output and the network is forced to converge to a specific final state (supervised learning). Each character is characterized by a vector of five components. For the formation of the network (multi-layer perceptron MLP), we started by a set of eighteen images, and we finished with 50 sets of 900 images, to find the best parameters that maximizes network (Figure.20).

Handwritten character sets	Numbers of characters	Manuscripts Test database	Printed Test database
1 set	18	62.43	60.20
10 sets	180	78.88	60.36
20 sets	360	81.57	61.86
30 sets	540	81.75	62.46
40 sets	720	84.23	63.83
50 sets	900	84.91	63.88
Numbers of images for test		2340 Images	3060 Images

Figure 16 : Experimental results for MLP

VII. HIDDEN MARKOV MODELS

Hidden Markov models (HMMs), (Choisy C., et al 2002), (Choisy C., et al 2002), (Choisy C., 2002) are statistical models parametric signal production, widely used in speech recognition and recognition of the writing later.

a) Markov Chain

A chain of discrete Markov of order n is a discrete stochastic process $X = \{X_t \mid t = 1, \dots, T\}$ with discrete random variables, checking the Markov property:

$$P(X_t = q_t \mid X_{t-1} = q_{t-1}, \dots, X_1 = q_1) = P(X_t = q_t \mid X_{t-1} = q_{t-1}, \dots, X_{t-n} = q_{t-n})$$

Ou $Q = \{q_1, \dots, q_n\}$ represents all the States.

b) Stationary chain

A Markov chain of order 1 is stationary if for any t and k there is:

$$P(X_t = q_i \mid X_{t-1} = q_j) = P(X_{t+k} = q_i \mid X_{t+k-1} = q_j)$$

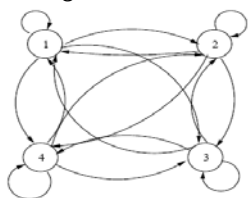
In this case, it defines a transition probability matrix $A = (a_{ij})$ such as:

$$a_{ij} = P(X_t = q_j \mid X_{t-1} = q_i)$$

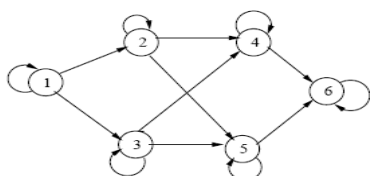
At a time given a any process.

c) Types of HMMs

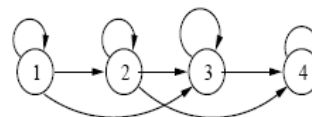
The main types of Markov models hidden are the ergodic and the right-left model.



Ergodic



left-right : parallel



left-right : sequential

d) Evaluation of the probability for observation:

Are $O = o_1 o_2 \dots o_T$ a suite of observations and $Q = q_1 q_2 \dots q_T$ suite of associated states. The probability for the observation of O , as the model β (or class) is equal to the sum on all possible States Q the joint probabilities of O and Q .

$$P(O|Y) = \sum P(O, Q|Y)$$

e) Calculation of $P(O|Y)$ using Forward-Backward function :

Observation can be done in two times:

- Emission of early observation O (1: t).
- Emission of the end of the observation O (t+1: t).

The evaluation of the observation is given by:

$$P(O|Y) = \sum \alpha(t, qi) * \beta(t, qi)$$

f) Recognition:

Knowing the class to which belongs the character it is compared to models λ_k , $k = 1, \dots, L$ of its class. The selected model will be the one to provide the best probability corresponding to the evaluation of its suite of primitive i.e: $\max (P(O / \lambda_k))$

With O : The suite of observation in this work is the vector of extraction.

λ_k : is the Markov model consisting of the transition matrix A , the observation matrix B and boot matrix Π .

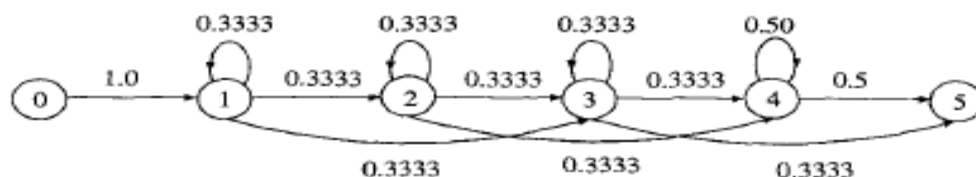


Figure 17 : Initial transitions

g) Chain of treatment

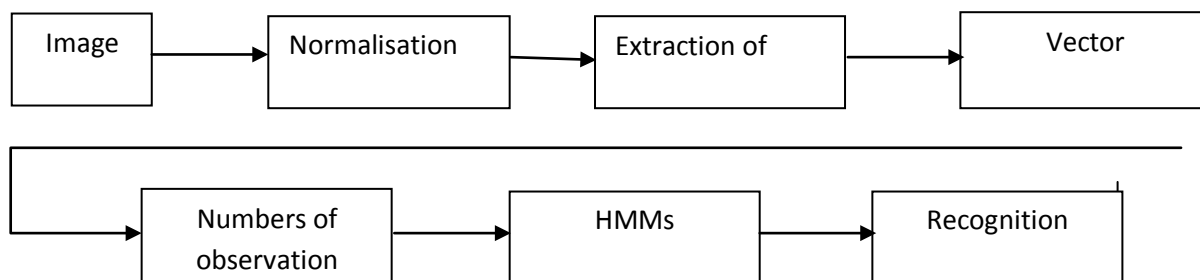


Figure 18 : The process for the HMM

h) Experimental results

For the classification with hidden Markov model, considering the values of the characteristic vectors obtained in recognition of the characters as a sequence of observations, it initializes the model and is sought with the Baum-Welch algorithm (learning) to find the

best probability that maximizes the parameters of the model (A: the transition matrix, B: the observation matrix, Π : the boot matrix). Each character is characterized by a vector of five components extraction. The experimental results are illustrated in the following figure (Figure 19).

Number of characters in the base	Type of characters used	Manuscripts Test database	Printed Test database
5400	18	92.22	77

Figure 19 : Experimental results for the HMM

VIII. HYBRID MODEL MLP/HMM

Under certain conditions, neural networks can be considered statistical classifiers by supplying output of a posteriori probabilities. Also, it is interesting to combine the respective capacities of the HMM and MLP for new efficient designs inspired by the two formalisms

a) The hybrid system

A perceptron can provide the probabilities of belonging $P(c_i/x(t))$ a vector model $x(t)$ to a class c_i . Several systems have been developed on this principle [22]. These systems have many advantages over approaches purely Markov. However, they are not simple to implement because of the number of parameters to adjust and the large amount of training data necessary to ensure the global model. In this section, we show how our hybrid system is designed. The architecture of the system consists of a multilayer perceptron upstream to a type HMM Bakis left-right. The hybrid system, including the figure 20 shows the overall

design scheme to provide probabilities of belonging to different classes. The system is composed of two modules: a neural module and a hidden Markov module.

MLP

HMM

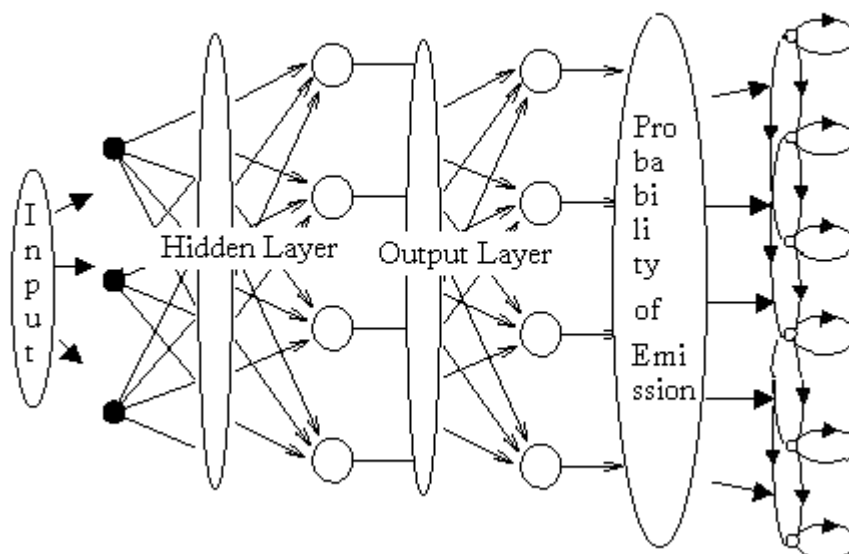


Figure 20 : Global schema design of the hybrid mode

The observations being the various classes, we associate with each State an observation. Also our goal is to find the best way of maximizing the probability of a sequence of observations, the method used for recognition with the hybrid model is the method used in

(section 7.6) for the hidden markov model. The HMM has in input two matrices and a vector: a matrix of transition probabilities, a matrix of emission probabilities which is the output of the MLP (posterior probabilities), and a vector representing the initial probability of States.

b) Results of the hybrid system

Number of characters in the base	Type of characters used	Manuscripts Test database	Printed Test database
5400	18	92.22	77

Figure 21: Experimental results for the HMM

Type of characters	MLP	HMM	MLP+HMM
Manuscripts	84.91	92.22	92.30
Printed	63.83	70	87.77

Figure 22 : Comparison between three methods of classification

IX. CONCLUSION

In this work, we use a method based on neural networks (multi-layer perceptron and back-propagation), the hidden Markov model and hybrid model MLP + HMM for the classification of manuscripts Tifinagh characters. A technique of extraction based on mathematical morphology is used in the phase of extraction of characteristics before the implementation in classification of characters. We prove that the HMM and MLP approaches can be a reliable method for the classification. We introduced the hybrid system to make more intelligent classification process by modeling the output of the neural network as probabilities of emission. When to initialize the settings of the hidden markov model. As perspectives we can increase the database and type characters used in recognition for the generalization of this method of extraction for the

recognition of all Tifinagh characters (33 characters according to IRCAM).

REFERENCES

1. Apurva A. Desai (2010) Gujarati handwritten numeral optical character reorganization through neural network. Pattern Recognition 43 (2010) 2582–2589.
2. Baidyk T., Kussul E. (2002), Application of neural classifier for flat image recognition in the process of microdevice assembly, Proceedings of the International Joint Conference on Neural Networks, Hawaii, USA 1 (2002) 160–164.
3. Ososkov G., (2003) Effective neural network approach to image recognition and control, Proceedings of International Conference on Physics and Control 1 (2003) 242–246.
4. Vitabile S., Gentile A., Sorbello F., (2002) A neural network based automatic road signs recognizer, Proceedings of the 2002 International Joint Conference on Neural Networks 3 2315–2320.
5. Serra J., 1982, Image Analysis and Mathematical Morphology. Academic Press, London.
6. Serra J., 1988, editor. Image Analysis and Mathematical Morphology.II : Theoretical Advances. Academic Press, London.

7. Soille P., 2003, Morphological Image Analysis : Principles and Applications. 2nd edition.
8. Dreyfus et. al (2001) : Réseaux de neurones. Méthodologie et applications. Eyrolles.
9. Bishop C. (1995) : Neural networks for pattern recognition. Clarendon Press - Oxford.
10. Haykin (1998) : Neural Networks. Prentice Hall.
11. Hertz, Krogh & Palmer (1991): Introduction to the theory of neural computation. Addison Wesley.
12. Thiria, Gascuel, Lechevallier & Canu (1997) : Statistiques et méthodes neuronales. Dunod.
13. Rosenblatt, F. (1962). Principles of Neurodynamics. Spartan.
14. Rumelhart, D. E., Hinton, McClelland, and Williams, R. J. (1986), Learning Internal Representations by Error Propagation.
15. Kussul E.M., Baidyk T.N., et al., (2001) Rosenblatt perceptrons for handwritten digit recognition, Proceedings of International Joint Conference on Neural Networks IJCNN 2 1516–1520.
16. Fletcher, R.(1987). Practical methods of optimization. New York : Wiley.
17. Ciarlet, P.G. (1989) Introduction to numerical linear algebra and optimisation. Cambridge : C.U.
18. Le Cun, Y., Boser B., Denker J., Henderson D., Howard R., Hubbard W., and Jackel L, (1990). Handwritten digit recognition with a back-propagation network. Advances in neural information processing systems. San Mateo, Morgan Kaufmann. 396-404.
19. LeCun Y., Bottou L, Bengio Y., Haffner P., (1998) Gradient-based learning applied to document recognition, Proceedings of the IEEE 86 (11) 2278–2344.
20. Jan A. Snyman (2005). Practical Mathematical Optimization : An Introduction to Basic Optimization Theory and Classical and New Gradient-Based Algorithms. Springer Publishing.
21. M. COHEN et al, "Combining neural networks and hidden Markov models for continuous speech recognition". ARPA Continuous Speech Recognition Workshop, Stanford University, September 21-22, 1992.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Ubiquitous Life Care Integrates Wireless Sensor Network And Cloud Computing With Security

By J.Jayashree, J.Vijayashree

VIT University, Vellore, TN, India

Abstract - The world we are living today is surrounded with different types of diseases so people started showing a great care for their health. This makes a rapid growth in healthcare systems. Ubiquitous life cares are attracted by people as they monitor patient health at home and at any time. In our proposed system Ubiquitous life care integrates wireless sensor networks with cloud computing and they also provide security to the data's on cloud. So it provides users with good care, reduces the cost and assures security.

GJCST Classification : C.2.1, D.4.6



Strictly as per the compliance and regulations of:



Ubiquitous Life Care Integrates Wireless Sensor Network And Cloud Computing With Security

J.Jayashree^α, J.Vijayashree^Ω

Abstract - The world we are living today is surrounded with different types of diseases so people started showing a great care for their health. This makes a rapid growth in healthcare systems. Ubiquitous life cares are attracted by people as they monitor patient health at home and at any time. In our proposed system Ubiquitous life care integrates wireless sensor networks with cloud computing and they also provide security to the data's on cloud. So it provides users with good care, reduces the cost and assures security.

I. INTRODUCTION

Now a day's people are showing more care and interest in their health and how to lead a healthy life with good life style. That is why there is a rapid increase in healthcare systems. Cost spend on healthcare is a big topic in all people in all countries.

Ubiquitous life care i.e. U Life Care are becoming more famous, interesting and attracted for researchers as they provide less cost with high quality at anywhere and at anytime. They monitor health at home. U Life Care are used for monitoring health of human and their activities and they share these details i.e. information among doctors. The quality and coverage they provide are really good.

Wireless sensor network is a emerging field and now becoming very attracted. they combines sensing, computation, and communication. now a days these sensors comes in many forms. they are very tiny so that these sensors can be inserted into a hand ring that these type of sensors are know as ring sensors and if they are inserted into belt they are know as belt sensors and these sensors are prepared with less cost. wireless sensors use Mesh networking protocol so that the connectivity they provide is really huge. The power of these sensor networks lies in the skill to organize large numbers of tiny nodes that assemble and configure themselves. In this project we use these sensors to monitor the health of patients.

Cloud computing technology has been derived from many different technologies like grid, virtualization and IT management. Cloud computing are becoming more concern and they will definitely modify the future. They have really attracted the clients. A small and medium enterprise gets benefited because there is no necessary to employ any special IT team. Cloud computing is concerned with resource sharing in distributed organization. There can be any number of

users i.e. participants who are really interested in participating in cloud and users are allowed to join and leave dynamically. These clouds offer many good benefits to the users. Cloud computing are considered very important as they provide dynamic resource sharing over internet. They also provide cloud computing infrastructure such as CPU capabilities, network and storage, they provide good infrastructure for many services and life care services are becoming very powerful through cloud computing. cloud computing provides life care with a high throughput and by reducing the cost. but the main drawback comes in the form of threat to privacy and security issues. There is also a fear of control loss of data. So security is now a very big issue in cloud computing. Through this paper, we present and pay attention to this problem and use some sort of security which are used to work for the protection of cloud i.e protecting the data or information on cloud.

II. RELATED WORKS

Author in [1] has focused on Human Activity Recognition Engine (HARE). HARE is a component which are used to provide real time data management, to human activities are detected and to manipulate the detected activities using ontologies. Here they have used Context information. HARE has been deployed for an Alzheimer's disease patient and they have made this with five different modules for activity recognition. HARE is made up of various components like Location Tracking, Activity Recognizer, Schema Mapping and XML Transformer and Activity Manipulation Engine which are used for tracking, recognizing, to transform the activities and to make decision. The experimental results were made and they have given achieved. But they have used this HARE only for Alzheimer's disease patients that is they have practiced only in this domain.

Author in [2][3] has presented an integration of wireless sensor networks and cloud computing for U-life care. Where the sensors monitor the health of the patient and the also the activities of the patient. These information are given to cloud and these information's are being shared by the doctors in cloud. The author has added some benefits by integrating the sensors and cloud computing. He has also explained about the existing system where the monitored data are not maintained in cloud instead in a centralizes server and that is considered to be unsafe. So here the

Author ^{α Ω} : P.G.Scholar; SITE: VIT University, Vellore – 632014, TN, India. E-mails : shakthi.shree@yahoo.com, shree.swathi@yahoo.com

authentication and access control are being used. Even here they have restricted this project to some domains only.

Author in [4] present detail about patient monitoring and the importance of that in current world. These monitoring are done under the cloud based telecare system. This type of system are really attracted by aged and people. People status are classified as static state and moving state. Strategy are classified as fully active, semi active and passive active. Medical Call Center(MCC) play a good role where it act as interface between the patient and medical staff and they are used for detecting the status of user continuously. the author has even added some of the clouds like hospital cloud and national cloud. The components related to this paper are E-house and mobile device and related software module. The added advantage to this is that they are applied to both fixed and mobile healthcare. The major drawback is that the author has not specified anything about the security of the data's.

III. PROPOSED SYSTEM

Lets consider a scenario where some XXX hospital has introduces a new medicine for some YYY disease and they have used this medicine to a ZZZ patient and this patient has been monitored continuously using wireless sensor which monitor the patient and sense his health condition and see whether that medicine is really working .the sensed data i.e the reports are being send to the hospital by uploading to the cloud.the hospital and the doctors doesn't want to

disclose about this medicine to the world.since these reports are now in the cloud some other hospital also have the chance of looking at these reports and even they can even try these to their patients .so in order to avoid these kind of problems ,we propose an idea where the reports being uploaded by the patient on the cloud and the same reports being downloaded by the hospital or the doctor that is both the patient and the doctor who want to access the reports has to be authenticated.

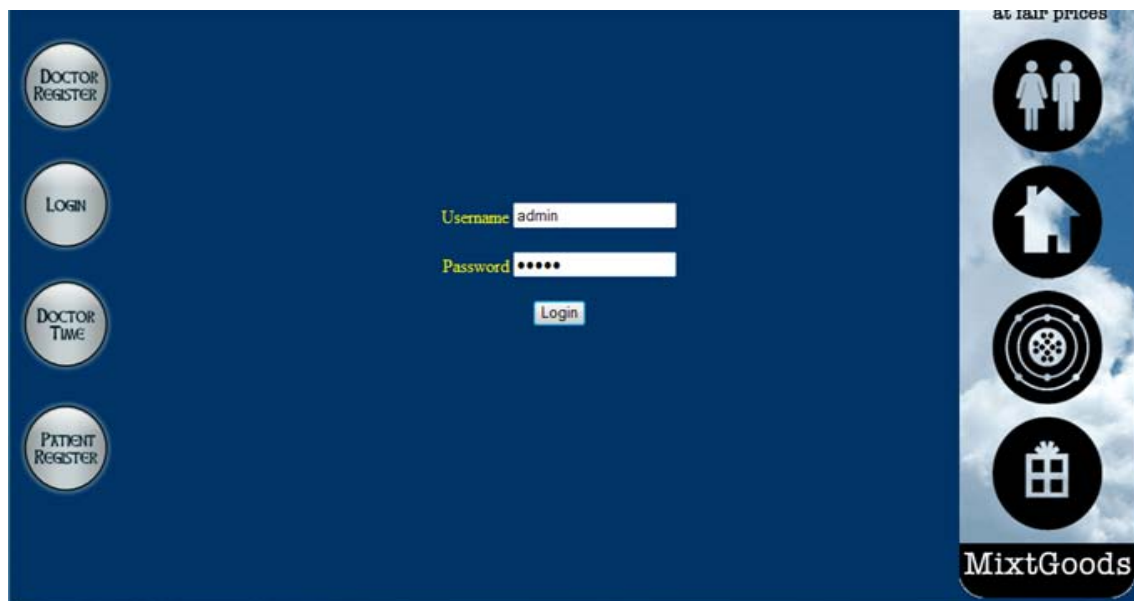
IV. METHODOLOGY

For proposing this idea the software requirements are Windows Operating System 7, JDK 1.6, JavaFX 1.3, XAMPP and MySQL.

A sensor captures all the activities of the patient and these datas or the reports are being transmitted to the cloud that is they upload these informations on cloud and if there is any noise data then they are filtered.now if the doctor want to access the data on cloud then he must be authenticate and permission must be granted so that he can access those reports being uploaded.

To access data on the Cloud, the user must authenticate and granted access permissions. When doctors, nurses want to access data, they must authenticate themselves first. After successful authentication, the Access Control module makes decision whether his/her access permission is allowed or not. If yes, it allows him/her to access to the Cloud data. Data is forwarded to authentic nurses and doctors.

V. RESULTS AND DISCUSSION



In fig.1 the admin login by using his username and password

Essential information,
dependable service

DOCTOR REGISTRATION

DOCTOR REGISTER

LOGOUT

DOCTOR TIME

Doctor Name: suresh

Password:

Gender: ☒ Male ☐ Female

Date of Birth: 11/5/1975

Address: chennai

Mobile: 9886532147

handmade goods by local artists at fair prices

Fig.2 shows Doctor registration. Doctor register by using his name, password, gender, date of birth, address and mobile number

at fair prices

DOCTOR REGISTER

LOGOUT

DOCTOR TIME

PATIENT REGISTER

Doctor Info

DoctorId	Specialist
doc4577	general
doc4978	eye
doc5992	eye
doc1345	general
doc3110	LS
doc9535	LAPROSCOPE
doc2320	eye
doc5737	general
doc7061	LS
doc0724	opopo
doc8195	Dental
doc1763	eye
doc6060	ORTHO

Time Alert

Doctor-ID: doc6060

☒ First ☐ Second ☐ Night

Okay

MixtGoods

Fig.3 shows Doctor Information in which the doctors id and his specification are listed.the doctors time shift is also given.



For Patients

PATIENT REGISTRATION

REGISTER
DOUBTS
LOGIN
DOCTOR AVAILABLE

handmade goods by local artists at fair prices

Patient Name: kali
 Password: ****
 Gender: ☒ Male ☐ Female
 Date of Birth: 4/5/1980
 Address: chennai
 Mobile: 9886532147
 Landline: 04452522
 Occupation: Engineer

Fig.4 here is the patient registration by having his name,password,gender,date of birth,address,contact number and his occupation.



Essential information, dependable service

DOCTOR REGISTRATION

DOCTOR REGISTER
LOGOUT
DOCTOR TIME

handmade goods by local artists at fair prices

Doctor Name: suresh
 Password: *****
 Gender: ☒ Male ☐ Female
 Date of Birth: 11/5/1975
 Address: chennai
 Mobile: 9886532147

Fig.5 shows Doctors administration



Fig.6 shows patient login through his username and password

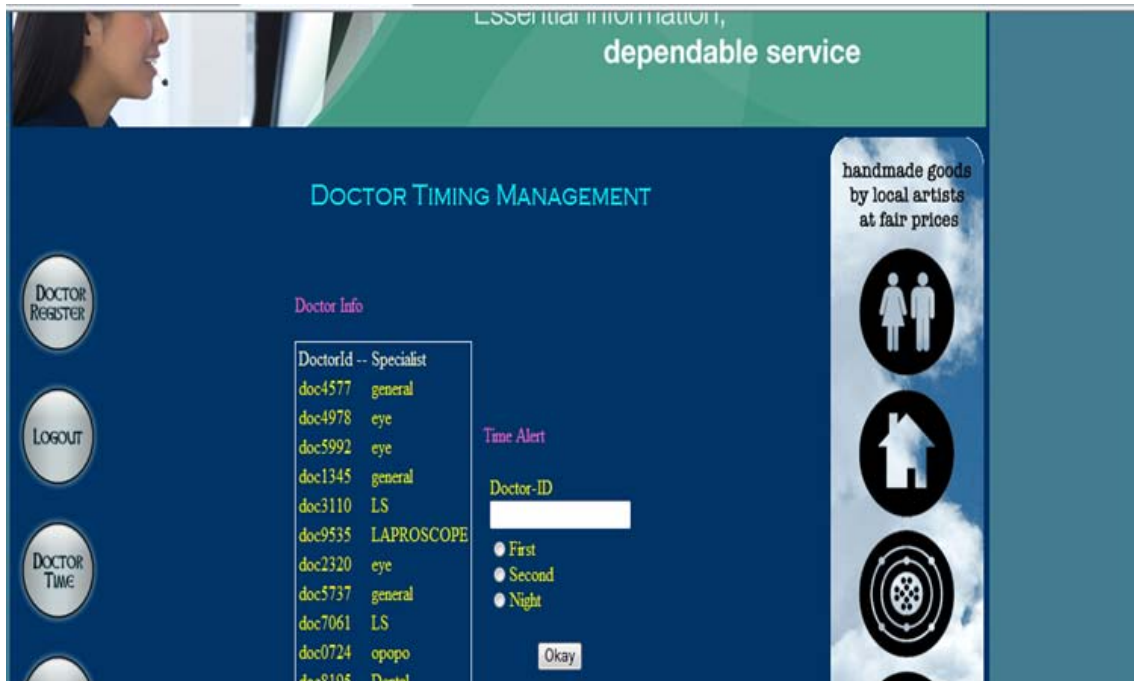


Fig.7 here is the Doctors time management which provides his id, specification and time



Fig.8 gives the details about the medicine information.it tells about the medicine,how much dosage to use and time of day.

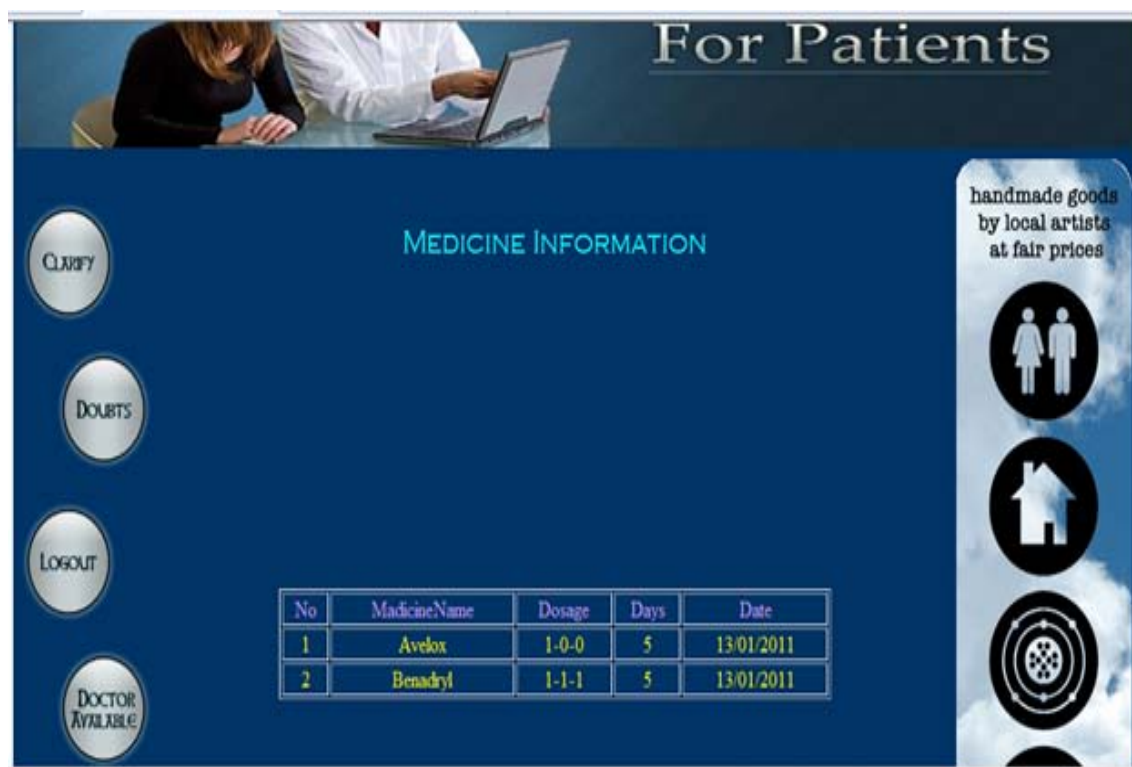


Fig.9 also provide some information about the medicine that is name of the medicine, dosage, for how many days they have to be used and date.

VI. CONCLUSION

Through this paper we have proposed that security can be provided to cloud which is integrated with sensor with U-life care. Both the patient and the doctor has to be authenticated before accessing the reports available on cloud. And in future new security measures can be introduces to make cloud more secured and this cloud security can be applied to many domains.

REFERENCES REFERENCES REFERENCIAS

1. Asad Masood Khattak, La The Vinh, Dang Viet Hung, Phan Tran Ho True, Le Xuan Hung, D. Guan, Zeeshan Pervez, Manhyung Han, Sungyoung Lee, Young-Koo Lee.: Context-aware Human Activity Recognition and Decision Making,IEEE,2010.
2. Xuan Hung Le, Sungyoung Lee¹, Phan Tran Ho Truc, La The Vinh, Asad Masood Khattak, Manhyung Han, Dang Viet Hung, Mohammad M. Hassan, Miso (Hyung-Il) Kim, Kyo-Ho Koo, Young-Koo Lee, Eui-Nam Huh.: Secured WSN-integrated Cloud Computing for u-Life Care,IEEE,2010.
3. H. Jameel, R.A. Shaikh, H. Lee and S. Lee.: Human Identification through Image Evaluation Using Secret Predicates. Topics in Cryptology - CT-RSA 07, LNCS 4377 (2007) 67–84.
4. Chen-Shie Ho and Kuo-Cheng Chiang.:Towards the Ubiquitous Healthcare by Integrating Active Monitoring and Intelligent Cloud.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Recognition of Tifinagh Characters Using Self Organizing Map And Fuzzy K-Nearest Neighbor

By S. Gounane, M. Fakir, B. Bouikhalene

FST SMS University Beni Mellal, Morocco

Abstract - In this paper we present a comparison between SOM (Self-Organization Map) neural network and Fuzzy K-nearest Neighbor algorithms and their application to handwriting Tifinagh character recognition. The Box approach proposed in [1] is used for features extraction. Experimentation is carried out on a limited database of nearly 200 samples. The results showed that Fuzzy K-Nearest Neighbor had a very good performance.

Keywords : *Tifinagh characters, Fuzzy k-Nearest Neighbor, Self Organizing Map, Fuzzy k-means, Features extraction.*

GJCST Classification : *I.2.3, F.1.1*



Strictly as per the compliance and regulations of:



Recognition of Tifinagh Characters Using Self Organizing Map And Fuzzy K-Nearest Neighbor

S. Gounane^α, M. Fakir^α, B. Bouikhalene^β

Abstract - In this paper we present a comparison between SOM (Self-Organization Map) neural network and Fuzzy K-nearest Neighbor algorithms and their application to handwriting Tifinagh character recognition. The Box approach proposed in [1] is used for features extraction. Experimentation is carried out on a limited database of nearly 200 samples. The results showed that Fuzzy K-Nearest Neighbor had a very good performance.

Keywords : Tifinagh characters, Fuzzy k - Nearest Neighbor, Self Organizing Map, Fuzzy k-means, Features Features extraction.

I. INTRODUCTION

The recognition of characters from scanned images of documents has been a problem that has received much attention in the fields of image processing, pattern recognition and artificial intelligence.

For many years, fuzzy logic and Artificial Neural Networks have been used in a wide range of problem domains: process control (where Fuzzy Controllers have been very popular), management and decision making, operations research, economics and pattern recognition and classification.

This paper presents an application of both SOM neural network and fuzzy k-Nearest Neighbor in recognition of handwritten Tifinagh characters. This paper is organized as follows: in section [1] a features extraction method, which is an essential step prior to pattern recognition, is described. Section [2] describes the architecture and the learning mechanism of the SOM neural networks. Section [3] presents the k-Nearest Neighbor algorithm. Section [4] gives results of the application of both SOM and Fuzzy k-NN on Tifinagh handwriting character recognition.

II. FEATURES EXTRACTION

Preprocessing techniques like thinning, slant correction and smoothening are applied. For extracting the features, the Box approach proposed in [1] is used here. This approach requires the spatial division of the

character image. The major advantage of this approach stems from its robustness to variation, ease of implementation and high recognition rate. Each character image is divided into $L \times H$ boxes so that the portions of character will be in some of these boxes. There could be boxes that are empty (Figure 1(a)). The choice of number of boxes is arrived at by experimentation. Elements of a Normalized Vector that describe the character is obtained by dividing the number of all black pixels present in this box by their total number for each box, (Figure 1 (b)).

One can easily see that this characterization is invariant of the character image dimensions. Hence an image of whatever size gets transformed into a vector of $L \times H$ predetermined dimensions.

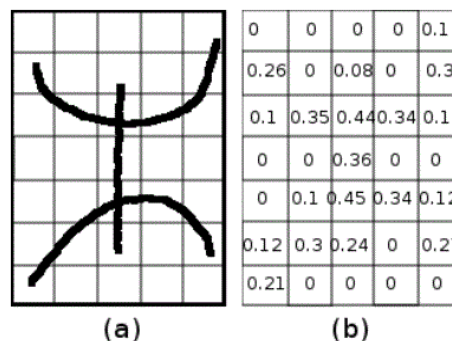


Figure 1 : Features extraction

III. SELF ORGANIZING MAP (SOM)

The Self-Organizing Maps, abbreviated SOM, was developed by Professor Teuvo Kohonen in the early 1980s. Self-organizing maps are a special class of artificial neural network, because those are based on competitive learning and the learning itself is unsupervised.

Also SOM is considered as a special case in data-mining, it can be applied to both clustering and projecting the data onto a lower dimensional display at the same time.

In a SOM network, there is an input layer and an output layer which is usually designed as two-dimensional arrangement of neurons that maps n dimensional input to two dimensional (Figure 2). It is basically a competitive network with the characteristic of self-organization providing a topology-preserving mapping from the input space to the clusters.

Author^α : Information Processing & Telecommunication Team, Dep. Of Computer Sciences – FST SMS University Beni Mellal, Morocco. E-mail : Gounane.said@gmail.com

Author^α : Information Processing & Telecommunication Team, Dep. Of Computer Sciences – FST SMS University Beni Mellal, Morocco. E-mail : fakfad@yahoo.fr

Author^β : Information Processing & Telecommunication Team – FP SMS University Beni Mellal, Morocco. E-mail : bouikhalene@yahoo.fr

The algorithm proceeds first by initializing the synaptic weights in the map for the neurons. This can be done by assigning them small values picked from a random number generator; in so doing no prior order is imposed on the feature map. Once the map has been properly initialized, there are three essential processes involved in the formation of the self-organizing map: Competition, cooperation and synaptic adaptation.

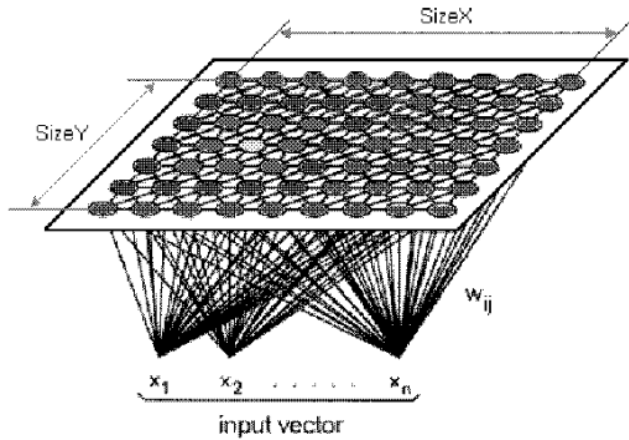


Figure 2 : SOM architecture

a) Competition

Let $X = [x_1, \dots, x_n]$ be the input pattern and $W_j = [w_{j1}, \dots, w_{jn}]$ the synaptic weight vector of each neuron j in the map. To find the best match of the input vector X with the synaptic weight vectors W_j , compare the inner products $W_j^T \cdot X$ for $j=1, 2, \dots, l$ (l is the number of output neurons) and select the largest denoted $i(X)$. Maximizing $W_j^T \cdot X$ is equivalent to minimizing $\|W_j - X\|$ [2] then one can have:

$$i(X) = \arg \min_j (W_j - X) \quad (1)$$

The neuron number $i(X)$ is called the winning neuron.

b) Cooperation

In neurobiology, a neuron that is firing tends to excite neurons in its immediate neighborhood more than those farther away. The same with output neurons in the SOM, a topological neighborhood around the winning neuron i is made and it decay smoothly with lateral distance. Another unique feature of the SOM algorithm is that the size of the topological neighborhood shrinks with time. Let $h_{j,i}(k)$ denote the topological neighborhood at time k , centered on the winning neuron i , and j denote a typical neuron of a set of excited (cooperating) neurons around winning neuron i . One can assume that the topological neighborhood $h_{j,i}$ is unimodal function of the lateral distance $d_{i,j}$, such that it satisfies:

1. $h_{j,i}$ is symmetric and attains its maximum value at

the winning neuron i for which $d_{j,i} = 0$.

2. The amplitude of $h_{j,i}$ decreases monotonically with increasing $d_{i,j}$, decaying to zero for $d_{i,j} \rightarrow \infty$.

A typical choice of $h_{j,i}$ is the Gaussian function

$$h_{j,i}(k) = e^{-\frac{d_{i,j}^2}{2\sigma^2(k)}} \quad (2)$$

$$\sigma(k) = \sigma_0 e^{-\frac{k}{\tau_1}} \quad (3)$$

Where τ_1 is the time constant of the algorithm.

c) Synaptic adaptation

By definition, for the network to be self-organizing (and unsupervised), the synaptic weight vector W_j of neuron j in the network is required to change in relation to the input vector X . This change is expressed by the equation as follows:

$$W_j(k+1) = W_j(k) + \alpha(k) h_{j,i(X)}(k) X(-W_j(k)) \quad (4)$$

Where $\alpha(k)$ is the learning-rate.

The exact form of learning-rate is not important. It can be linear, exponential or inversely proportional. However it should be time varying. In particular, it should start at an initial value α_0 , and then decrease gradually with increasing time n . This requirement can be satisfied by using an exponential learning-rate, as shown by:

$$\alpha(k) = \alpha_0 e^{-\frac{k}{\tau_2}} \quad (5)$$

Where τ_2 is another time constant of the SOM algorithm.

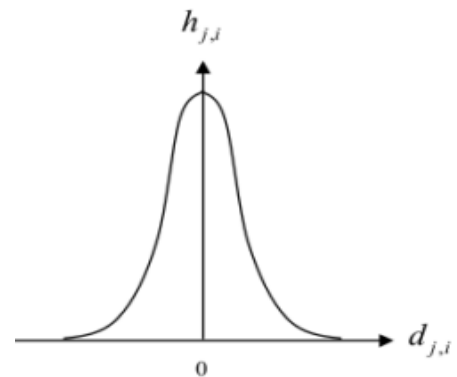


Figure 3 : example of topological neighbourhoods (Gaussian)

IV. FUZZY K-NEAREST NEIGHBOR

Fuzzy kNN is part of supervised learning that has been used in many applications in the field of data mining, statistical pattern recognition, image processing

and many others. Some successful applications are including recognition of handwriting and satellite image. Fuzzy K nearest neighbor algorithm is very simple. It works based on an unsupervised clustering algorithm like fuzzy k-means algorithm to determine prototypes representing clusters. Then the minimum distance from the query instance to these prototypes is used to determine the K-nearest neighbors. After and basing on membership function we take the neighbor with the maximum membership to be the prediction of the query instance.

a) Fuzzy k-means

The Fuzzy k-means clustering algorithm is an improvement of the k-means algorithms. K-Means are the simplest methods of clustering data. The fuzzy K-means algorithm uses a set of N unlabeled feature vectors and classifies them into k classes, where k is given by the user.

From these N feature vectors, k of them are randomly selected as initial seeds. The feature vectors are assigned to the closest seeds depending on its distance from it and on a given membership function. The mean of features belonging to a class is taken as the new center. The features are reassigned; this process is repeated until convergence.

In a fuzzy clustering, a pattern vector can belong to all clusters with different degrees given by a membership function [figure 4]. One can prove that such a function exists [5], and for each cluster C_i it is defined as follows:

$$f_i(X) = \frac{1/d^2(X, c_i)^{\frac{1}{m-1}}}{\sum_{i=1}^N 1/d^2(X, c_i)^{\frac{1}{m-1}}} \quad m \in \mathbb{R}, \quad m \geq 1 \quad (6)$$

Where c_i is the center of the class C_i , and The parameter m is defined by the user to adjust the fuzziness of the clustering. The hard clustering case is obtained by taking $m = 1$.

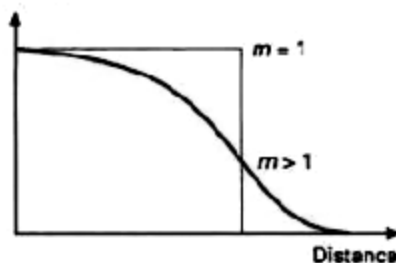


Figure 4 : Membership function

b) Fuzzy K-means algorithm

1. Choose randomly the k prototypes c_i .
2. Compute the degree of membership of all feature vectors X_j in all clusters C_i : $\mu_{ij} = f_i(X_j)$.
3. Compute new cluster prototypes as follows:

$$c_i = \frac{\sum_{j=1}^N (\mu_{ij})^m x_j}{\sum_{j=1}^N (\mu_{ij})^m} \quad (7)$$

4. Iterate back and force between (2) and (3) until the memberships or cluster centers for successive iteration differ by more than some prescribed value ϵ .

c) Fuzzy K-Nearest Neighbor

The Fuzzy k-NN classifiers consist on proximity measures. They are ideally suited for modeling the non parametric distribution on handwritten word recognition data.

For a given character X, the fuzzy classifier computes the membership of X in different classes $C_1, C_2, \dots, C_N$. The character X is allocated to the class for which the membership function yields the maximum value. By a learning process (Fuzzy k-mean, k-mean, LVQ ...) we assign a number of prototypes P_j for each class C_i . After generating the k-Nearest Neighbor prototypes P_j for a character (by distance similarity), the degree of membership of the vector X to the class C_i : $\mu_i(X)$, can be calculated as follow:

$$\mu_i(X) = \frac{\sum_{j=1}^k \mu_{ij} / d^2(X, P_j)^{\frac{1}{m-1}}}{\sum_{j=1}^N 1/d^2(X, P_j)^{\frac{1}{m-1}}} \quad (8)$$

Where μ_{ij} is the membership of the prototype P_j to the class C_i . To compute this value tow methods are proposed:

Using a crisp assignment of P_j to C_j :

$$\mu_{ij} = \begin{cases} 1 & P_j \in C_i \\ 0 & P_j \notin C_i \end{cases} \quad (9)$$

Or by using the k-NN to determine k-nearest neighbor of P_j , then the number of neighbor belonging to the same class as P_j denoted n_j is used to compute μ_{ij} as follows:

$$\mu_{ij} = \begin{cases} 0.51 + 0.49n_j / k & j = i \\ 0.49n_j / k & j \neq i \end{cases} \quad (10)$$

V. RESULTS

A java application was developed in order to test and compare those two algorithms. To extract features each character was divided into boxes, the input layer of the SOM neural network was made of 63 neurons and 35 for the output layer (number of character to recognize), for the fuzzy k-nearest neighbor the fuzziness factor m is equal to 2 and number of neighbors' k is equal to 3.

As there is no standard database available for handwritten Tifinagh characters, a database was created from samples of different writing styles with different sizes. The database also includes some complex samples that are even hard to recognize by careful inspection Figure 5: Samples used for training (figure 5). This database contains 200 samples (about 5 samples per character) divided into two disjoint sets, one for training both SOM and Fuzzy k-NN, and another for testing these two algorithms.

The features extraction method used here is scaling invariant, that's why there is always a recognition problem between Tifinagh character $\circ$ and $\circ$, where $\circ$'s are mistaken as $\circ$'s and vice versa. Some other recognition problems are with and and , $\circ$ and $\circ$ but not persistent as the first one.

Experiments shows that the best results are given by the Fuzzy k-NN where in many cases it success to recognize handwritten characters that SOM algorithm fails to (figure 6).

	Sample1	Sample2	Sample3	Sample4
				
				
				
				
				
				
				

Figure 5 : Samples used for training

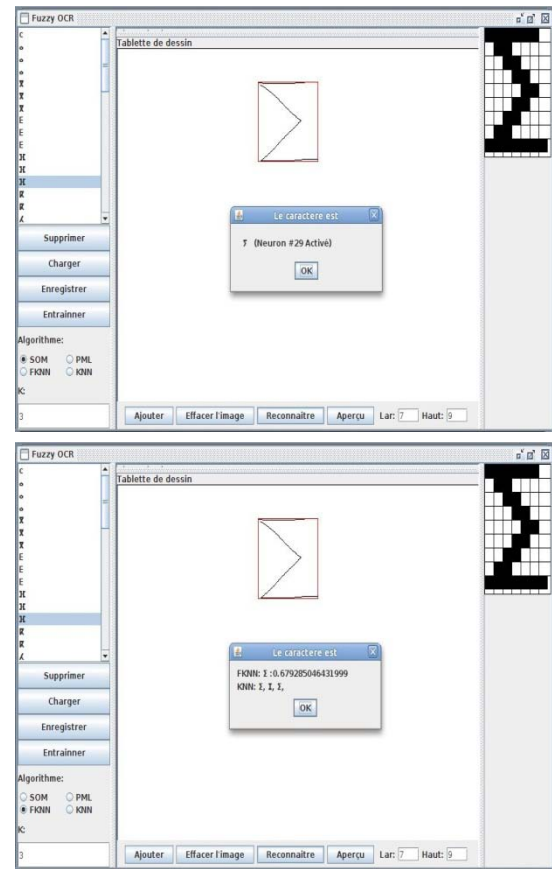


Figure 6 : handwritten character. SOM algorithm (left) fails () but not Fuzzy k-NN(right) ()

VI. CONCLUSION

In this paper two clustering algorithms are presented, Self Organizing Map neural network and the Fuzzy k-Nearest Neighbor. We applied these two algorithms to Tifinagh character recognition. The box approach was used to extract features from the character image. Fuzzy k-NN was more robust and fast than the SOM.

REFERENCES REFERENCES REFERENCIAS

1. M.Hanmandlu, A.V.Nath, A.C.Mishra, & V.K.Madasu. (2007). Fuzzy Model Based Recognition of Handwritten Hindi Numerals using bacterial foraging. 6th IEEE/ACIS ICIS.
2. Haykin, & Simon. (1999). Neural Networks: a comprehensive fondation. New jersey: Prentice-hall inc.
3. A.Mingoti, S., & Lima, J. O. (2005). Comparing SOM neural network with Fuzzy c-means, K-means and traditional hierarchical clustering algorithms. European Journal of Operational Research.
4. D.Slezak, B.Kang-Junzhong, T.Kim, & G.Kuroda. Segnal processing, Image processing and Pattern recognition. New York: Springer.
5. J.Sagau. (n.d.). Logique flou en classification. Techniques de l'Ingénieur , H3638.

6. (J.C.), & BEZDEK. (1981). Pattern recognition with fuzzy objective function algorithms. Plenum Press.
7. A.Tellaeche, Burgos-Artizzu, X. P., G.Pajares, & A.Ribeiro. (2007). On Combining Support Vector Machines and Fuzzy K-Means in Vision-based Precision Agriculture. World Academy of Science, Engineering and Technology.
8. F.Zheng. (n.d.). K-Means-based fuzzy classifier.
9. Nasser, S., Alkhaldi, R., & Vert, G. (n.d.). A Modified Fuzzy K-means Clustering using Expectation Maximization.
10. Nassery, P., & Faez, K. (n.d.). Signature pattern recognition using psudo Zernike moments and fuzzy logic classifier.
11. O.Matan, R.K.Kiang, C.E.Stenard, & B.Boser. (1990). Handwritten character recognition using neural network architectures. 4th USPS advanced Technology Conference, Washington D.C, 1003-1011.
12. S.Nasser, R.Alkhaldi, & G.Vert. (n.d.). A Modified Fuzzy K-means clustering using expectation Maximization.
13. S.Theodoridis, & Koutroumbas, K. (2009). Pattern recognition (fort edition). Elsevier.
14. V.Ganapathy, & Liew, K. L. (2008). Handwritten Character Recognition Using Multiscale Neural.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Framework for Costing Service-Oriented Architecture (SOA) Projects Using Work Breakdown Structure (WBS) Approach

By Yusuf Lateef Oladimeji, Olusegun Folorunso, Akinwale Adio Taofeek,
Adejumobi, A. I

University of Agriculture, Abeokuta, Ogun State, Nigeria

Abstract - With end users demanding faster response time and management demanding lower costs and more flexibility, Service Oriented Architecture (SOA) projects are becoming more complex and brittle. Proper costing and identification of feasible benefits of SOA projects are quickly becoming a significant influence in the mainstream of all industries. SOA is intended to improve software interoperability by exposing dynamic applications as services. Current SOA quality metrics pay little attention to service complexity as an important key design feature that impacts other internal SOA quality attributes. Due to this complexity of SOA, cost and effort estimation for SOA-based software development is more difficult than that of traditional software development. Unfortunately, there is little or no effort about cost and effort estimation for SOA-based software. Traditional software cost estimation approaches are inadequate to address the complex service-oriented systems. Although numerous sources expound on the technical advantages of SOA as well as listing praises for their intuitive and qualitative benefits, until now no one has provided a reliable and quantifiable result from SOA implementations currently in production. This paper proposes a novel framework based on Work Breakdown Structure (WBS) approach for cost estimation of SOA-based software by dealing separately with service parts. The WBS framework can help organizations simplify and regulate SOA implementation cost estimation by explicit identification of SOA-specific tasks in the WBS. Furthermore, both cost estimation modelling and software sizing work can be satisfied respectively by switching the corresponding metrics within this framework. We provide an example case study to demonstrate proposed metrics and we also investigate the benefit of SOA to its adopters.

Keywords : *Service-Oriented Architecture (SOA), Software Cost Estimation, Work Breakdown Structure (WBS), Framework, Return on Investment (ROI).*

GJCST Classification : *D.2.9, D.2.8*



Strictly as per the compliance and regulations of:



© 2011. Yusuf Lateef Oladimeji, Olusegun Folorunso, Akinwale Adio Taofeek, Adejumobi, A. I. This is a research/review paper, distributed under the terms of the Creative Commons Attribution-Noncommercial 3.0 Unported License (<http://creativecommons.org/licenses/by-nc/3.0/>), permitting all non commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

A Framework for Costing Service-Oriented Architecture (SOA) Projects Using Work Breakdown Structure (WBS) Approach

Yusuf, Lateef Oladimeji^α, Olusegun Folorunso^α, Akinwale, Adio Taofeek^β, Adejumobi, A. I^ψ

Abstract - With end users demanding faster response time and management demanding lower costs and more flexibility, Service Oriented Architecture (SOA) projects are becoming more complex and brittle. Proper costing and identification of feasible benefits of SOA projects are quickly becoming a significant influence in the mainstream of all industries. SOA is intended to improve software interoperability by exposing dynamic applications as services. Current SOA quality metrics pay little attention to service complexity as an important key design feature that impacts other internal SOA quality attributes. Due to this complexity of SOA, cost and effort estimation for SOA-based software development is more difficult than that of traditional software development. Unfortunately, there is little or no effort about cost and effort estimation for SOA-based software. Traditional software cost estimation approaches are inadequate to address the complex service-oriented systems. Although numerous sources expound on the technical advantages of SOA as well as listing praises for their intuitive and qualitative benefits, until now no one has provided a reliable and quantifiable result from SOA implementations currently in production. This paper proposes a novel framework based on Work Breakdown Structure (WBS) approach for cost estimation of SOA-based software by dealing separately with service parts. The WBS framework can help organizations simplify and regulate SOA implementation cost estimation by explicit identification of SOA-specific tasks in the WBS. Furthermore, both cost estimation modelling and software sizing work can be satisfied respectively by switching the corresponding metrics within this framework. We provide an example case study to demonstrate proposed metrics and we also investigate the benefit of SOA to its adopters.

Keywords : *Service-Oriented Architecture (SOA), Software Cost Estimation, Work Breakdown Structure (WBS), Framework, Return on Investment (ROI).*

Author ^α : Department of Computer Science University of Agriculture, Abeokuta, Ogun State, Nigeria GSM Phone no : +234 803 403 7098; E-mail : truevisionconsulting@yahoo.com

Author ^α : Department of Computer Science University of Agriculture, Abeokuta, Ogun State, Nigeria GSM Phone no : +234 803 564 0707; E-mail : folorunsolusegun@yahoo.com

Author ^β : Department of Computer Science University of Agriculture Abeokuta, Ogun State, Nigeria GSM Phone no : +234 806 286 7612; E-mail : atakinwale@yahoo.com

Author ^ψ : Department of Electrical and Electronics Engineering University of Agriculture Abeokuta, Ogun State, Nigeria GSM Phone no : +234 703 321 5455; E-mail : engradejumobi@yahoo.com

I. INTRODUCTION

A well-developed understanding of the Return on Investment (ROI) for SOA has been a complex undertaking [1]. This is due in part to the deficiency in comprehensive historical data on which to base any such model. For the most part, SOA often exist as pilot projects than as full-blown production systems, and even those rare production-quality systems that do exist are too new for use in understanding critical issues to the ROI equations, such as reusability and redeployment. Most organizations that want to build an SOA don't have a clue on how to approach the cost estimation process. In many cases, they grossly underestimate the cost of their SOA, hoping the management won't notice, this is done to get approval and reveal the higher costs later after investment may have been made and too late to go back. This is not a good management practice. The other problem militating against building a comprehensive cost model for SOA is the need to separate the service cost that results in SOA from any well-designed application and the specific or incremental cost that obtains from a well designed SOA application built on services architecture. Software cost estimation for Service-Oriented Architecture (SOA) development confronts more challenges than for traditional software development. One of the main reasons is the architectural difference in SOA compared to traditional software development. Josuttis [2] has pointed out that distributed processing would be inevitably more complicated than non-distributed processing, and any form of loose coupling will increase complexity. Meanwhile, the more complexity involved in a system, the more difficulty the designers or engineers have to understand the implementation process and thus the system itself [3]. In other words, people have to devote more effort to accurate manipulations when performing more complicated tasks. In practice, building a true heterogeneous SOA for a wide range of operating environments may take years of development time if the company does not have sufficient SOA experience and expertise [4]. It is difficult to foresee and justify the cost and effort of developing an SOA application before the project starts. The problem of SOA cost estimation has not been addressed adequately in the existing literature.

The current cost estimation approaches for traditional software development are inadequate for complex service-oriented software. For example, COCOMO II cannot arrive at global cost approximation for the entire SOA application development, and expert judgment may easily fall into traps of uncertainty or bias because of the complexity of the SOA. This paper proposes a novel framework by employing a Work Breakdown Structure (WBS) approach in an attempt to deal with cost estimation problem for SOA based software development. Within this WBS framework, services are classified into three primitive types and one combined type according to different development processes. Cost estimation for developing primitive services can be handled as sub-problems that are small and independent enough to be solved. For combined services, the division procedure will emerge recursively until all the resulting separated services are primitive. The cost and effort of service integration is then calculated gradually following the reverse division sequence. The application of the WBS cost estimation framework is demonstrated using a case study. The result shows that the proposed framework can simplify and regulate the complicated development cost estimation for SOA-based applications. The business goals and objectives of SOA are to increase agility and reduce costs while the technical goals and objectives are to increase usability, improve maintainability and reduce redundancy [5]. SOA hold out the promise for a brave new world of applications development, deployment, and reuse that many proponents believe will usher in unprecedented levels of Return on Investment (ROI) for a domain that has long suffered from cost overruns and excessive, often unjustified expenditures. The ability to lower the cost of integration while improving the leveragability of key software and business process assets are only a few of the reasons why the ROI of service-oriented architectures and composite applications is thought to herald a new economic reality for IT and business development. Ease of use and lower training costs, lower cost of deployment, faster time to market, improved business requirement matching, and better multi-channel deployment are among the myriad reasons the technologies are so eagerly awaited by business and IT managers alike.

II. RELATED WORK

a) SOA Services

SOA is a collection of services with well-defined interfaces and a shared communications model. A service is a coarse-grained, discoverable, and self contained software entity that interacts with applications and other services through a loosely coupled, often asynchronous, message-based communication model [6]. A system or application is designed and implemented to make use of these services. This

developed capability may itself provide services within the overall SOA. The underlying idea of SOA is that it would be cheaper and faster to build or modify applications by composing them out of limited-purpose components that can communicate with each other because the components strictly adhere to interface rules [7]. The advent of the Internet and World Wide Web (WWW) introduced a new wave of research on collaborative product development environment [8][9][10][11][12]. Yusuf et al., [13] Observed that the Internet is no longer a simple network of computers but a network of potential services in which the functional views of services need to be clearly defined during the design of an Internet-based distributed engineering system. The most common form of SOA is that of Web services in which all of the following apply: service interfaces are described using Web Services Description Language (WSDL), payload is transmitted using Simple Object Access Protocol (SOAP) over Hypertext Transfer Protocol (HTTP), and Universal Description, Discovery and Integration (UDDI) is optionally used as the directory service [14]. However, WSDL, SOAP, and HTTP are not the only foundation on which an SOA can be built. Other technologies such as CORBA and IBM's Web sphere can be used as part of the messaging backbone of an SOA.

b) Work Breakdown Structure

A Work Breakdown Structure (WBS) is a hierarchical decomposition (tree structure) of the work required to accomplish a goal. It is developed by starting with the end objective and successively re-dividing it into manageable components in terms of size, duration, and responsibility [15]. However, it is often done as a modification of an existing WBS for a similar project. It is an essential starting input to both estimation and to scheduling. In essence, it provides the chart of accounts for a project. To know what something cost, it needs to exist as a task in the WBS. In large projects, the approach is quite complex and can be as much as five or six levels deep. Usually, items at the same level of hierarchy are in the order they are executed, although this is not required. Traditionally, definition of the WBS is left to vendors with the integrated master schedule and price proposal based upon it included as part of the RFP response. More often than not, the organization of the WBS in software development follows the traditional "waterfall" method of system development. The primary constraint is that the WBS fulfils the requirements of the statement of work [16]. Since the development of the software within services is about the same as traditional development, we suggest Breakdown of SOA into Services from a WBS perspective.

c) COCOMO II

COCOMO II (Constructive Cost Model) [17] is one of the best-known and best-documented algorithmic models, which allows organizations to

estimate cost, effort, and schedule when planning new software development activities. Tansey and Stroulia [18] have attempted to use COCOMO II to estimate the cost of creating and migrating services. They reported that COCOMO II should be extended to accommodate new characteristics of SOA based development. COCOMO II is generally inadequate to accommodate the cost estimation needs for SOA-based software development. When considering the declarative composition specifications, a fundamentally different development process may be adopted in SOA-based software. Based on the Internet technologies, SOA-based software can be realized as a composition of loosely coupled services with well-defined interfaces and consistent communication protocols. These services hide technical details, and are not restricted to any specific technology. In other words, the service implementation is programming language and platform independent. Therefore, an SOA-based application could comprise the combination of all possible development strategies and development processes. Consequently, although the COCOMO II model has a large number of coefficients such as effort multipliers and scale factors, it is difficult to directly justify the cost estimation for SOA-based software development. On the other hand, considering the difference between component orientation and service orientation [19], the COCOMO II model by itself is inadequate to estimate effort required when reusing service-oriented resources. COCOMO II considers two types of reused components, namely black-box components and white-box components. Black-box components can be reused without knowing the detailed code or making any change to it, while white-box components have to be modified with new code or integrated with other reused components before it can be reused. Similarly, within the SOA framework, there are black-box services that can be adopted directly, and white-box services that should be ported from legacy systems. Nevertheless, taking black-box reuse for instance, the difference between code-level and service-level reuse is significant. Whether a code-level component is suitable or not for reuse should be understood and revealed by using reverse engineering or reengineering [20] according to the real situation. Comparatively, the contractually reusable and loosely coupled service can be reused directly through service discovery techniques, for example semantic annotation and quality of service.

d) Function Point Analysis and Software Sizing

Size prediction for the constructed deliverables has been identified as one of the key elements in any software project estimation. SLOC (Source Line of Code) and Function Point are the two predominant sizing measures. Function Point measures software system size through quantifying the amount of functionality provided to the user in terms of the number

of inputs, outputs, inquires, and files. In practice, Function Point can be used continuously throughout the entire software development life cycle, which provides the essential value of what the software is and what it does with data from the user's viewpoint. Santillo attempts to use the Function Point method to measure software size in an SOA environment [21]. After comparing the effect of adopting the first and second generation methods, that is the International Function Point User's Group (IFPUG) and Common Software Measurement International Consortium (COSMIC) respectively, Santillo identifies several critical issues. The prominent one is that SOA is functionally different from traditional software architectures, because the "function" of a service should represent a real-world self-contained business activity [2]. More issues appear when applying IFPUG to software system size measurement. For example, the effort of wrapping legacy code and data to work as services cannot be assigned to any functional size. Measuring with the COSMIC approach, on the contrary, is supposed to satisfy the typical sizing aspects of SOA-based software. However, there is a lack of guidelines for practical application of COSMIC measurement in SOA context. In addition to the application of Function Points, Liu et al. [22] use Service Points to measure the size of SOA-based software. The software size estimation is based on the sum of the sizes of each service.

$$Size = \sum_{i=1}^n (P_i \times P)$$

Where P_i is an infrastructure factor with empirical value that is related to the supporting infrastructure, technology and governance processes; P represents a single specific service's estimated size that varies with different service types, including existing service, service built from existing resources, and service built from scratch. This approach implies that the size of a service-oriented application depends significantly on the service type. However, the calculation of P for various services is not discussed in detail.

e) SMART and SMAT-AUS Framework

The "Service-Oriented Migration and Reuse Technique" (SMART) was developed to assist organizations in analyzing legacy capabilities for use as services in an SOA. SMART was derived from the Options Analysis for Reengineering (OAR) method developed at the SEI that was successfully used to support analysis of reuse potential for legacy components [23]. SMART gathers a wide range of information about legacy components, the target SOA, and potential services to produce a service migration strategy as its primary product. However, SMART also produces other outputs that are useful to an organization whether or not it decides on migration.

Information-gathering activities are directed by the Service Migration Interview Guide (SMIG). The SMIG contains questions that directly address the gap between the existing and target architecture, design, and code, as well as questions concerning issues that must be addressed in service migration efforts. Use of the SMIG assures broad and consistent coverage of the factors that influence the cost, effort, and risk involved in migration to services. Unlike SMART, SMAT-AUS [24] is a framework that is developed to determine the scope and estimate cost and effort for SOA projects. This framework reveals not only technical dimension but also social, cultural, and organizational dimensions of SOA implementation. When applying the SMAT-AUS framework to SOA-based software development, Service Mining, Service Development, Service Integration and SOA Application Development are classified as separate SOA project types. For each SOA project type, a set of methods, templates and cost models and functions are used to support the cost and effort estimation work for each project time which are then used to generate the overall cost of an SOA project (a combination of one or more of the project types). Except for the SMART (Software Engineering Institute's Service Migration and Reuse Technique) method [25] that can be adopted for service mining cost estimation, currently there are no other metrics suitable for the different projects beneath the SMAT-AUS framework. Instead, some abstract cost-estimation-discussions related to aforementioned project types can be found through a literature review. Umar and Zordan [26] warn that both gradual and sudden migration would be expensive and risky so that costs and benefits must be carefully weighed. Bosworth [27] gives a full consideration about complexity and cost when developing Web services. Liu et al. [22] directly suggest that traditional methods can be used to estimate the cost of building services from scratch. Since utilizing solutions based on interoperable services is part of service-oriented integration (SOI) and results in an SOI structure, Erl [28] gives a bottom line of effort and cost estimation for cross-application integration: "The cost and effort of cross-application integration is significantly lowered when applications being integrated are SOA-compliant." A generic SOA application could be sophisticated. But this can be handled in SMAT-AUS by breaking the problem into more manageable pieces (i.e. a combination of project types) however specifying how all of these pieces are estimated and the procedure required for practical estimation of software development cost for SOA-based systems is still being developed.

f) ROI of SOA Based on Traditional Component Reuse

Barry Boehm provided two useful formulas when estimating the cost of software systems reuse. One formula is from the provider's point of view, while

the other is from consumer's [29]; Provider-focused formula:

Relative Cost of Writing for Reuse (RCWR) = Cost of Developing Reusable Asset / Cost of Developing Single-User Asset

Consumer's formula : Relative Cost of Reuse (RCR) = Cost of Reuse Asset / Cost of Develop Asset from Scratch

Poulin Jeffery [30] examined large-scale SOA service providers to estimate the value ranges for these formulas in practice. His data shows that RCWR ranges between 1.15 and 2.0 with median of 1.2, while RCR ranges between 0.15 and 0.80 with a median of 0.50. In other words, Paulin work suggests that creating reusable software component for a broad audience takes more resources (15% to 100% more) than creating a less generic point solution. The 20% of the total cost of development directed towards reuse, a factor Poulin calls Relative Cost of Reuse (RCR) would represent an impressive number in the pre-object-oriented development world, but in the world of service oriented architectures and component application; it is believed that 80% is a more accurate figure. This ability to reuse the majority of the software development by an organization is one of the key attributes of SOA development, and while the number will vary greatly from one development organization to another, it is our believe that the early adopters will see this or an even greater degree of reuse simply because initial SOA development will target precisely those applications and business processes that have the greater reuse potential. Mili et al., [31] has published a variety of RCWR factor values that have been developed since the early 1990's based on the experiences of a number of sources. Discussion about cost estimation for SOA implementation also appears in industry. Linthicum [32] outlines some general guidelines for estimating the cost of an SOA application. According to these guidelines, the calculation of SOA cost can be expressed as a sum of several cost analysis procedures.

Cost of SOA = (Cost of Data Complexity
+ Cost of Service Complexity
+ Cost of Process Complexity
+ Enabling Technology Solution)

Furthermore, Linthicum also provides some detailed specification. For example, the basic element Complexity of the Data Storage Technology is figured as a percentage between 0% and 100% (Relational is 30%, Object-Oriented is 60%, and ISAM is 80%). Nevertheless, the other aspects of the calculation are suggested to follow similar means without clarifying essential matters. Meanwhile, Linthicum reminds that the notable problem is that this approach is not a real metric. Additionally, SOA based software is inevitably more complicated than traditional software [2]. It is therefore doubtful that Data Complexity, System

Complexity, Service Complexity and Process Complexity are sufficient to represent the complexity of SOA-based systems. As shown, both academia and industry have published little work relating to estimating costs for SOA-based software. In particular, there is not a solution to satisfy the development cost estimation for SOA-based software. We attempt to address these issues by providing a SOA cost estimation framework in this paper.

III. METHODOLOGY

a) SOA-Based Software Cost Estimation Using Work Breakdown Structure (WBS) Approach

WBS approach is a "Division of labour" or "Divide and conquered" method which can be traced back to as early as 200BC [33], when the Babylonian reciprocal table of Inakibit-Anu was used to facilitate searching and sorting numerical values. However, the first description of the divide and conquered algorithm appears in John Mauchly's article discussing its application in computer sorting [33]. Nowadays, the approach is applied widely in areas such as Parallel Computing [34], Clustering Computing [35], Granular Computing [36], and Huge Data Mining [37]. The principle underlying WBS is shown in Figure 1. That is to recursively decompose SOA into sub problems (services) until all the sub-problems are sufficiently simple enough, and then to solve the sub-problems (cost the services). Resulting solutions (costs) are then recomposed to form an overall solution. Adopting this principle will lead to different subroutines for different sub-problems. Normally, some or all of the sub-problems are of the same type as the input problem, thus WBS procedure can be naturally expressed recursively. The QuickSort [33] algorithm has such procedure.

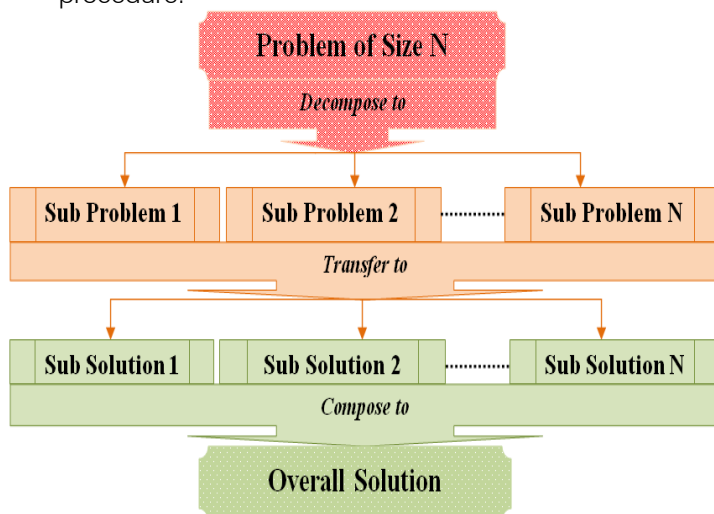


Figure 1: Principle of Work Breakdown Structure (WBS)

The advantages of applying WBS approach to SOA problems are numerous, and can be classified as followings:

- *Structural Simplicity* : Profiting from perhaps the simplest structuring technique, WBS is a high priority strategy to resolve problems not only in the SOA field but also in computing generally, politics and sociology fields. No matter where the approach is applied the solution structure can be expressed explicitly in a program-like function such as: Solution(x) is equivalent to:

```
IF IsBase(x)
Then SolveDirectly(x)
Else Compose(Solution(Decompose(x)))
```

Where x is the original problem that will be solved through *Solution* procedure, *IsBase* is used to verify whether the problem x is primitive or not, which returns TRUE if x is a basic problem unit, or FALSE otherwise. *Solve Directly* presents the conquer procedure. *Decompose* is referred to as the decomposing operation, while *Compose* is referred to as the composing operation

- *Computational Efficiency* : WBS can be used for designing fast algorithms. In appropriate application scenarios, the approach leads to asymptotically optimal cost for solving the problems. A problem of size N can be broken into a bounded number P of sub-problems of size N/P step by step, and all the basic sub-problems have constant-bounded size. Then the algorithm will have $O(\text{Mog}N)$ worst-case program execution performance. Normally, the consequence is more flexible because the size and the number of tasks can be decided at run-time.
- *Parallelism* : Since sub-problems in the individual division stage are logically and physically independent, the WBS approach can be naturally executed in parallel procedures. For computing problems, it is suitable for application in parallel machines due not only to the independent problem grains but also the efficient use of cache and deep memory hierarchies [38]. In fact, it has been considered as one of the well-known parallel programming paradigms.
- *Capability of Solving Complexity* : Through Dismantle of an overall goal into smaller and independent sub-problems, the SS strategy can provide adaptation scalability and variability, and can be used in the areas of engineering to reduce and manage complexity. Those complicated cases, such as resolutions for conceptually difficult problems, and approximate algorithms for NP-hard problems, are usually based on the divide and conquer principle. Given these merits, WBS can be considered a suitable and effective approach to accommodate complex problems such as cost estimation for SOA-based software development, where individual measures must be carried out independently. The following sections discuss its applications in SOA cost estimation.

b) Service Classification

Implementing SOA could be complex and onerous, while complexity measurement for SOA-based system is still an open question [39]. Chaos [40] claims that the complexity is restricting some SOA implementations. For the same reason, there are also many challenges to estimate the cost and effort of SOA-based software development. Fortunately, the major advantages of SOA are mainly reusability and composability with an emphasis on extensibility and flexibility, at a high level of granularity and abstraction. In other words, SOA-based software can be naturally divided into a set of loosely coupled services. These services can then be classified through their different features. Krafzig et al. [41] has identified that distinguishing services into classes is extremely helpful when properly estimating the implementation and maintenance cost, and the cost factors may vary depending on the service type. However, there is no standard way to categorize services. Service classification can be different for different purposes, for example differentiating services according to their target audience [2], categorizing services through their business roles and responsibilities [28], and classifying services by using their background techniques and protocols [42]. Services in our work are characterized as follows:

- *Available Service* (basic service type), when the service already existing i.e. it may be provided by a third party or inherited from legacy SOA based systems.
- *Migrated Service* (basic service type), is the service to be generated through modifying or wrapping reusable traditional software component(s).
- *New Service* (basic service type), is the service to be developed from scratch.
- *Combined Service* is the service arising from the combination of any above three types of basic services or other combined services.

Through this type of classification, four different development areas are identified in SOA projects. These areas present both a decomposition process that results in Service Discovery, Service Migration, and Service Development, and a recomposition process that is Service Integration. The cost estimation for overall SOA-based software development can then be separated into these smaller areas with corresponding metrics. Therefore, the WBS approach is a feasible attempt for SOA-based software cost estimation following this development oriented service classification.

c) WBS Cost Estimation Frameworks

The proposed cost estimation framework for SOA-based software follows the WBS principle. Firstly, through the service-oriented analysis, the SOA project is divided into basic services recursively. Secondly,

different sets of metrics are adopted to satisfy the cost and effort estimation for different service development processes. The total cost and effort of the SOA project will be calculated through the service integration procedure as shown in Figure 2.

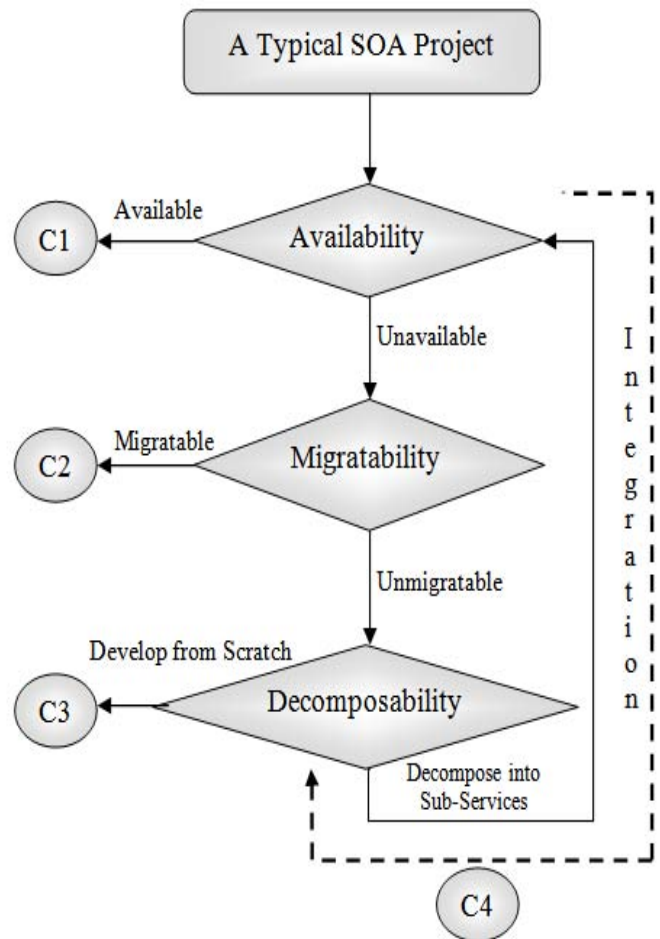


Figure 2: Procedure of SOA Project Development Cost Estimation based on WBS.

C1 is the cost estimation model or software size measurement used to accomplish modelling or sizing work for discovering available services, C2 represents migrating potential services, C3 represents developing new services, and C4 is the cost estimation model or size measurement for calculating the service integration effort. The Decomposability condition depends on the design and real situations whether the current service should be further divided or developed as a whole. The framework in Figure 2 presents the generic process of SOA cost estimation using the WBS method. To precisely describe the WBS based cost estimation for SOA-based software development, the complete process was expressed in pseudo code (Table 1). We define the stage that service division occurs as the service levels, and the combined service stands in a higher level next to its successive component services.

Table 1: Algorithm of SOA Project Development Cost Estimation Based on WBS

```
//Treat the project at the highest-level service S to be analyzed.
double SoaCostEstimation(service S) {
double cost = 0;
//Determine the type of S according to the design and real situations.
switch (the type of S) {
case AVAILABLE:
cost += The cost of service discovery;
break;
case MIGRATABLE:
cost += The cost of service migration (service wrapping);
break;
case NEW:
cost += The cost of service development;
break;
default:
//Divide S into component services at lower level.
for each component service in S
cost += SoaCostEstimation(component service);
cost += The cost of service integration for component services in S;
break;
}
return cost;
}
```

The SOA project itself is treated at the highest-level coarse-grain service, which is also the initial input parameter of SoaCostEstimation function. Within the body of SoaCostEstimation function, the cost of the input service development will be estimated directly if the service belongs to those three basic types, or recursively calculated by analyzing and composing the cost and effort of the development for component services. When composing individual service development costs into the overall SOA-based software development cost, the strategy of supposed service integration is progressed level-by-level instead of integrating the services all at once. The reason of adopting such a strategy is that, according to our work, service integration occurring in different levels will make different contributions to the total cost and effort of the project development. A real example can be used to show the application process of the WBS based cost estimation framework for SOA-based software development in practice, which is demonstrated in the next section.

IV. IMPLEMENTATION

We employ Visualization RCD Beam for Service Oriented Architecture (VisRCDBeam for SOA) implemented by Yusuf et al [43] as an application case study. There are two reasons for choosing this case: The VisRCDBeam for SOA case study characterizes all the service types listed in the previous section, and there are a limited number of services which are adequate for illustrative purpose in this paper.

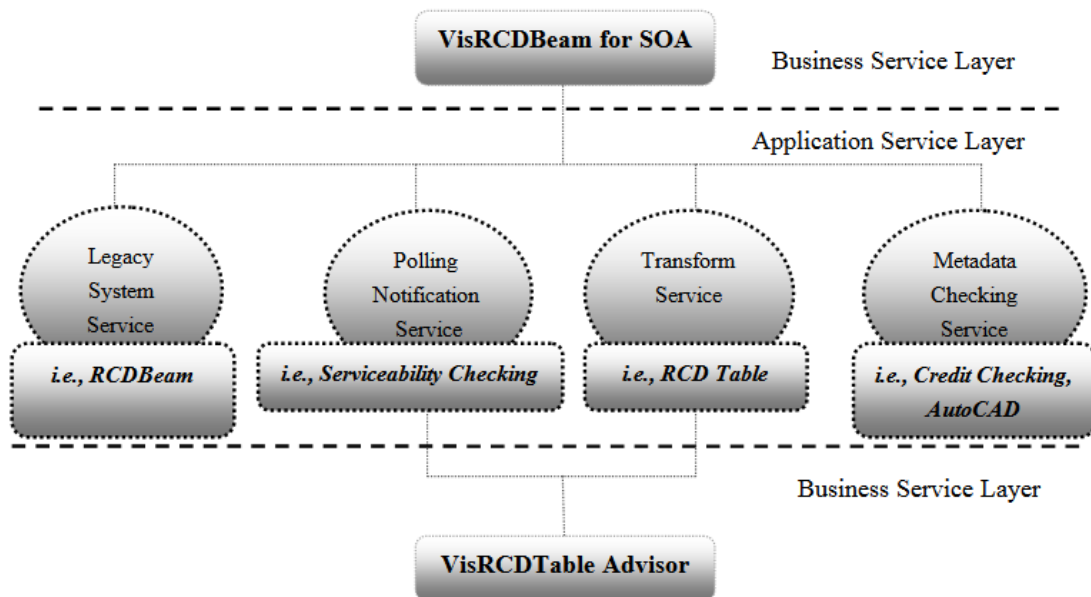


Figure 3 : Redesigned Automation System of VRCD SOA

VisRCDBeam for SOA is a SOA tool for the analysis and design of Reinforced Concrete Structures. To improve the working efficiency of Reinforced Concrete Designers, a service-oriented analysis was

conducted, which decomposed the business process logic into series candidates. The tool revealed the requirements of two business services in higher level and four application services in lower level. The

improved automation system is represented in Figure 3 following current disciplines:

- RCDBeam interface is the Legacy System Service which is migrated from the previous project.
- Serviceability checking represents the Polling Notification Service which is a coarse grain service containing check for minimum and maximum area of steel and check for deflection functional services. The Transform Service is the RCD table for picking bar sizes. These are new services that were developed from scratch.
- Credit checking for authentication and security, and AutoCAD interface for visualization purposes is represented by the Metadata Checking Service. They are the available service provided by third party.
- "VisRCDBeam for SOA" Service and VisRCDDTable Advisor Service are both combined services containing all or some of above basic services. The procedure of cost and effort estimation for developing this redesigned service-oriented project is illustrated in Figure 4. The detailed steps are elaborated as follows:

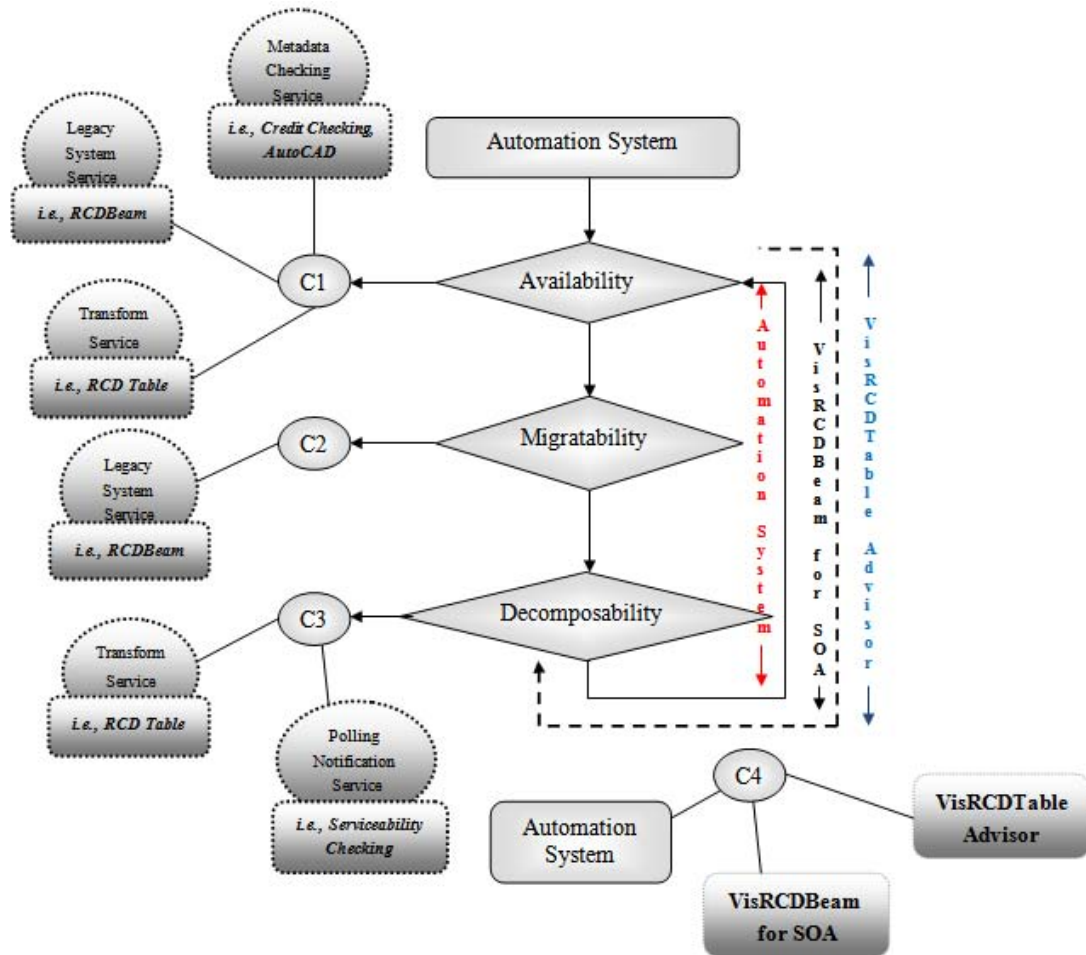


Figure 4 : Procedure of Cost and Effort Estimation for "VisRCDBeam for SOA" case study

- Divide the Automation System into *VisRCDBeam for SOA* Service and *VisRCDDTable Advisor* Service.
- Divide the *VisRCDBeam for SOA* Service into its four basic component services.
- Estimate the cost and effort of discovering the available Metadata Checking Service (i.e., *Credit checking and AutoCAD services*) by using corresponding metrics C1.
- Estimate the cost and effort of migrating the Legacy System Service (i.e., *RCDBeam*) by using corresponding metrics C2.
- Estimate the cost and effort of developing the Polling Notification Service (i.e., *Serviceability* checking) and Transform Service (i.e., *RCD Table*) by using corresponding metrics C3.
- Estimate the cost and effort of integrating the above four component services into the *VisRCDBeam for SOA* Service by using corresponding metrics C4.
- Divide the *VisRCDDTable Advisor* Service into its two basic component services.
- Notice that Legacy System Service (i.e., *RCDBeam*) and Transform Service (i.e., *RCD Table*) have both been taken into account.
- Estimate the cost and effort of mining the Legacy System Service (i.e., *RCDBeam*) and Transform Service (i.e., *RCD Table*) by using corresponding

metrics C1. Since these two services are in the same project and can be directly identified, the cost and effort here can be treated as zero in this special case.

- j. Estimate the cost and effort of integrating the above two component services into the *VisRCDDTable Advisor* Service by using the corresponding metrics C4.
- k. Estimate the cost and effort of integrating the *VisRCDDBeam for SOA* Service and *VisRCDDTable Advisor* Service into the Automation System by using the corresponding metrics C4.
- l. Sum up all the estimation results to calculate the total cost and effort of the Automation System development. Through the demonstration of the *VisRCDDBeam for SOA* case, the WBS framework is proven helpful for simplifying and regulating the SOA-based software cost estimation. Moreover, all the simplified cost estimation problems are independent enough to be solved in parallel. The uniform and explicit working procedure within this WBS framework is then a feasible attempt to SOA based software cost estimation.

V. SOA BENEFITS

SOA benefit organizations in different ways, depending on the respective goals and the manner in which SOA is applied. We have generalized the list of common benefits and certainly not exhaustive. It is merely an indication of the potential this architectural platform has to offer.

a) Improved integration (and intrinsic interoperability)

SOA can result in the creation of solutions that consist of inherently interoperable services. The net result is intrinsic interoperability, which turns a cross-application integration project into less of a custom development effort, and more of a modeling exercise. The cost and effort of cross-application integration is significantly lowered when applications being integrated are SOA-compliant.

b) Inherent reuse

Service-orientation promotes the design of services that are inherently reusable. Building services to be inherently reusable results in a moderately increased development effort and requires the use of design standards. Subsequently leveraging reuse within services lowers the cost and effort of building service-oriented solutions.

c) Streamlined architectures and solutions

The concept of composition is another fundamental part of SOA. It is not, however, limited to the assembly of service collections into aggregate services. The WS platform is based in its entirety on the principle of composability. This aspect of service-oriented architecture can lead to highly optimized

automation environments, where only the technologies required actually become part of the architecture. Realizing this benefit requires adherence to design standards that govern allowable extensions within each application environment. Benefits of streamlined solutions and architectures include the potential for reduced processing overhead and reduced skill-set requirements (because technical resources require only the knowledge of a given application, service, or service extension).

d) Leveraging the legacy investment

The industry-wide acceptance of the Web services technology set has spawned a large adapter market, enabling many legacy environments to participate in service-oriented integration architectures. This allows IT departments to work toward a state of federation, where previously isolated environments now can interoperate without requiring the development of expensive and sometimes fragile point-to-point integration channels. Though still riddled with risks relating mostly to how legacy back-ends must cope with increased usage volumes, the ability to use what we already have with service-oriented solutions that we are building now and in the future is extremely attractive. The cost and effort of integrating legacy and contemporary solutions is lowered. The need for legacy systems to be replaced is potentially lessened.

e) Establishing standardized XML data representation

On its most fundamental level, SOA is built upon and driven by XML. As a result, an adoption of SOA leads to the opportunity to fully leverage the XML data representation platform. A standardized data representation format (once fully established) can reduce the underlying complexity of all affected application environments. Past efforts to standardize XML technologies have resulted in limited success, as XML was either incorporated in an ad-hoc manner or on an "as required" basis. These approaches severely inhibited the potential benefits XML could introduce to an organization. With contemporary SOA, establishing XML data representation architecture becomes a necessity, providing organizations the opportunity to achieve their goal, the cost and effort of application development is reduced after a proliferation of standardized XML data representation is achieved.

f) Focused investment on communications infrastructure

Because Web services establish a common communications framework, SOA can centralize inter-application and intra-application communication as part of standard IT infrastructure. This allows organizations to evolve enterprise-wide infrastructure by investing in a single technology set responsible for communication. The cost of scaling communications infrastructure is reduced, as only one communications technology is required to support the federated part of the enterprise.

g) "Best-of-breed" alternatives

Some of the harshest criticisms laid against IT departments are related to the restrictions imposed by a given technology platform on its ability to fulfill the automation requirements of an organization's business areas. This can be due to the expense and effort required to realize the requested automation, or it may be the result of limitations inherent within the technology itself. Either way, IT departments are frequently required to push back and limit or even reject requests to alter or expand upon existing automation solutions. SOA won't solve these problems entirely, but it is expected to increase empowerment of both business and IT communities. A key feature of service-oriented enterprise environments is the support of "best-of-breed" technology. Because SOA establishes a vendor-neutral communications framework, it frees IT departments from being chained to a single proprietary development and/or middleware platform. For any given piece of automation that can expose an adequate service interface, we now have a choice as to how we want to build the service that implements it. The potential scope of business requirement fulfillment increases, as does the quality of business automation.

h) Organizational agility

Agility is a quality inherent in just about any aspect of the enterprise. A simple algorithm, a software component, a solution, a platform, a process all of these parts contain a measure of agility related to how they are constructed, positioned, and leveraged. How building blocks such as these can be realized and maintained within existing financial and cultural constraints ultimately determines the agility of the organization as a whole. Much of service-orientation is based on the assumption that what you build today will evolve over time. One of the primary benefits of a well-designed SOA is to protect organizations from the impact of this evolution. When accommodating change becomes the norm in distributed solution design, qualities such as reuse and interoperability become commonplace. The predictability of these qualities within the enterprise leads to a reliable level of organizational agility. However, all of this is only attainable through proper design and standardization. Change can be disruptive, expensive, and potentially damaging to inflexible IT environments. Building automation solutions and supporting infrastructure with the anticipation of change seems to make a great deal of sense. A standardized technical environment comprised of loosely coupled, composable, and interoperable and potentially reusable services establishes a more adaptive automation environment that empowers IT departments to more easily adjust to change. Further, by abstracting business logic and technology into specialized service layers, SOA can establish a loosely coupled relationship between these two enterprise

domains. This allows each domain to evolve independently and adapt to changes imposed by the other, as required. Regardless of what parts of service-oriented environments are leveraged, the increased agility with which IT can respond to business process or technology-related changes is significant. The cost and effort to respond and adapt to business or technology-related change is reduced.

VI. DISCUSSION

The aim of *visRCDBeam for SOA* is to upgrade its automation system so that it could remain competitive with other RCD tools and continue its business relationship with its primary client. We proceeded with a service-oriented analysis that decomposed its business process logic into a series of service candidates. This revealed the need for the following potential services and service layers: (a) A business service layer consisting of two tasks centric business services namely *visRCDBeam for SOA* and *VisRCDDTable Advisor*, (b) An application service layer comprised of four application services. Each business process was represented with a task-centric business service that would act as a controller for a layer of application services. Reusability and extensibility in particular were emphasized during the design of the application services. We intend to have the initial SOA to consist of services that supported both of its current business processes, while being sufficiently extensible to accommodate future requirements without too much impact. To realize the *visRCDBeam for SOA* tool, we compose these services into a two-level hierarchy where the parents *VisRCDBeam for SOA* and *VisRCDDTable advisor* business services coordinate the execution of all application services. Unlike many of the current cost estimation approaches, the proposed WBS framework uses a set of metrics to satisfy the development cost estimation for SOA-based software. The WBS concentrates on the software development process. It list Service Discovery as an individual cost estimation area as well as Service Migration, Service Development, and Service Integration. The framework estimates overall cost and effort through the independent estimation activities in four different development areas of an SOA application. WBS framework is generic and flexible by switching different types of metrics; it could satisfy different requirements of SOA-based software cost estimation such as building cost estimation model, measuring software size, and predicting the overall cost ultimately. One issue is that there are currently few available metrics for the detailed cost estimation for SOA-based software development. Future research should develop new metrics to resolve this issue. Meanwhile, some reusable existing metrics can be integrated into the proposed WBS framework, for example Tansey and Stroulia's work [18] [related to

Service Development and SMART method [26] are related to Service Migration. Over all, instead of trying to enumerate SOA project types, the WBS framework unifies and regulates the cost and effort estimation for SOA-based software development.

VII. CONCLUSION

Poor project management could bring failure to SOA. Gone are the days when one person is the SOA architect, developer, data architect, network architect and security specialist. The complexity of SOA should not be underestimated. Failure to implement and adhere to SOA governance is an imperative issue; the development effort is shifting from building services to consuming services. Vendors could be allowed to drive the architecture but relying too much on vendors can be a disaster. Software cost estimation plays a vital role in software development projects, especially for SOA-based software development. Current cost estimation approaches for SOA-based software are inadequate due to the architectural difference and the complexity of SOA applications. This paper offers a WBS cost estimation framework for SOA-based software development. Based on the principle of Divide and Conquer theory, this framework can be helpful for simplifying the complexity of SOA cost estimation. By hosting different sets of metrics, this generic framework will be suitable not only for the complete cost estimation work but also for the partial requirements, such as building estimation model, and measuring the size of SOA applications. We have fulfilled our original goals by producing proper costing of an SOA project that supports two service-oriented solutions. Online transaction is now possible. New requirements can be accommodated with minimal impact. The standard application service layer will likely continue to offer reusability functionalities to accommodate the fulfilment of new requirements. And any functional gaps will likely be addressed by extending the services without significantly disrupting existing implementations. Furthermore, should we decide to replace our task-centric business services with an orchestration service layer in the future, the abstraction established by the existing application service layer will protect the application services from having to undergo redevelopment. We have established a legacy system service (which is essentially a wrapper service for graphics drawing) as part of its application service layer it has opened up a generic and point that can facilitate integration. There is an old saying that you cannot manage what you cannot measure. By increasing the number of "moving pieces" in IT solutions, SOA increases the number of pieces that require measurement. Given the relative immaturity of the SOA paradigm, it is particularly important now, when best practices have not yet been established and the

understanding of cause and effect is limited. Indeed, the inability to collect cost and schedule data at the task level may be part of the reason why so many case studies in SOA only present project-level estimates of averted cost.

REFERENCES REFERENCES REFERENCIAS

1. Cresswell, Anthony M., (2004). Return on Investment In Information Technology: A Guide for Managers, Center for Technology in Government, University at Albany, August 2004.
2. Josuttis, N. M. (2007). SOA in Practice: The Art of Distributed System Design, Sebastopol: O'Reilly Media, Inc.
3. Cardoso, J. (2005). "How to Measure the Control-Flow Complexity of Web Processes and Workflows," Workflow Handbook 2005, Layna Fischer, Apr. 2005, pp. 199-212.
4. Jamil, E. (2009). "SOA in Asynchronous Many-to One Heterogeneous Bi-Directional Data Synchronization for Mission Critical Applications," We Do Web Sphere, Jul. 2009. [Online]. Available: <http://wedowebosphere.de/news/1528/SOA%20in%20Asynchronous%20Many-toone%20Heterogeneous%20BiDirectional%20Data%20Synchronization%20>. [Accessed: Nov. 2009].
5. Robert D. Bryan, Brand K. Niemann and Kartik Mecheri (2011). Perspectives on SOA ROI. Karsun Solutions Enterprise Modernization Expert. Geoff Raines mitre.org/news/digests/enterprise_modernization/09_08/external.html.
6. Brown, A; Johnston, S., and Kelly, K. (2002). *Using Service-Oriented Architecture and Component-Based Development to Build Web Service Applications*. Santa Clara, CA: Rational Software Corporation.
7. Lloyd Brodsky (2010). Meaningful Cost-Benefit Analysis for Service-Oriented Architecture Projects. Proceedings of the Seventh Annual Acquisition Research Symposium Thursday Sessions Volume II, Monterey, California, U.S. Government or Federal Rights License.
8. Maxfield, J., Fernando, T. and Dew, P. (1995). A Distributed Virtual Environment for Concurrent Engineering", in IEEE Proceedings on Virtual Reality Annual International Symposium, March 1-15, Research Triangle Park, North Carolina, pp.162-170.
9. Dong, A and Agogino, A.M. (1998). Managing design information in enterprise-wide CAD using 'smart drawings', Computer-Aided Design, 30 (6), 425-435.
10. Roy, U. and Kodkani, S.S. (1999). Product modelling within the framework of the World Wide Web, IIE Transactions, 31 (7), 667-677.
11. Huang, G.Q. and Mak, K.L. (2000). WeBid: A Web-based Framework to Support Early Supplier



- Involvement in New Product Development, Robotics and Computer Integrated Manufacturing, 16 (2-3), 169-179.
12. Kan, H.Y., Duffy, V.G. and Su, C.J. (2001). An Internet Virtual Reality Collaborative Environment for Effective Product Design, Computers In Industry, 45 (2), 197-213.
13. Yusuf Lateef Oladimeji, Olusegun Folorunso, Akinwale Adio Taofeek and Adejumbi, I.A. (2011). "Service Oriented Application in Agent Based Virtual Knowledge Community". Computer and Information Science Journal, Vol. 4, No. 2; March 2011.
14. Lewis, Grace and Wrage, Lutz. (2011). *Approaches to Constructive Interoperability* (CMU/SEI-2004-TR-020). Pittsburgh, PA: Software Engineering Institute, Carnegie Mellon University, 2005. <http://www.sei.cmu.edu/publications/documents/04.reports/04tr020.html>.
15. Senn, James A. (1987). *Information Systems in Management*. 3rd ed. Belmont, California: Wadsworth Publishing.
16. Mishan, E. J. (1976). *Cost-Benefit Analysis*. 2nd ed. New York: Praeger Publishers.
17. Boehm, B.W., Abts, C., Brown, A.W., Chulani, S., Clark, B.K., Horowitz, E. R., Madachy, D.J., Reifer, and Steece, B. (2000). Software Cost Estimation with COCOMO II. New Jersey: Prentice Hall PTR, Aug. 2000.
18. Tansey, B. and Stroulia, E. (2007). "Valuating Software Service Development: Integrating COCOMO II and Real Options Theory," Proc. the First International Workshop on the Economics of Software and Computation, IEEE Press, May 2007, pp. 8-8, doi: 10.1109/ESC.2007.11.
19. Stojanovic Z. and Dahanayake, A. (2005). Service-Oriented Software System Engineering: Challenges and Practices. Hershey, PA: IGI Global, Apr. 2005.
20. Sommerville, I. (2006). Software Engineering, 8th ed.. London: Addison Wesley, Jun. 2006.
21. Santillo, L. (2007). "Seizing and Sizing SOA Application with COSMIC Function Points," Proc. the 4th Software Measurement European Forum, Rome, Italy, May 2007.
22. Liu, J., Xu, Z., Qiao, J., and Lin, S. (2009). "A Defect Prediction Model for Software Based on Service Oriented Architecture using EXPERT COCOMO," Proc. Chinese Control and Decision Conference (CCDC '09), IEEE Press, Jun. 2009, pp. 2591-2594, doi: 10.1109/CCDC.2009.5191800.
23. Bergey, J., O'Brien, L., and Smith, D. (2002). "Using the Options Analysis for Reengineering (OAR) Method for Mining Components for a Product Line," 316-327. *Software Product Lines: Proceedings of the Second Software Product Line Conference (SPLC2)*. San Diego, CA, August 19-22, 2002. Berlin, Germany: Springer, 2002.
24. O'Brien, L. (2009). "A Framework for Scope, Cost and Effort Estimation for Service Oriented Architecture (SOA) Projects," Proc. 20th Australian Software Engineering Conference (ASWEC'09), IEEE Press, Apr. 2009, pp. 101-110, doi: 10.1109/ASWEC.2009.35.
25. Lewis, G., Morris, E., O'Brien, L., Smith, D., and Wrage, L. (2005). "SMART: The Service-Oriented Migration and Reuse Technique," CMU/SEI-2005-TN-029, Software Engineering Institute, USA, Sept. 2005.
26. Umar, A. and Zordan, A. (2009). "Reengineering for Service Oriented Architectures: A Strategic Decision Model for Integration versus Migration," Journal of Systems and Software, vol. 82, Mar. 2009, pp. 448-462, doi: 10.1016/j.jss.2008.07.047.
27. Bosworth, A. (2001). "Developing Web Services," Proc. 17th International Conference on Data Engineering (ICDE 2001), IEEE Press, Apr. 2001, pp. 477-481, doi: 10.1109/ICDE. 2001. 914861.
28. Erl, T. (2005). Service-Oriented Architecture: Concepts, Technology, and Design. Crawfordsville: Prentice Hall PTR, Aug. 2005.
29. Progress Actional (2008). "Web Services and Reuse" <http://www.actional.com/resources/whitepapers/SOA-Worst-Practices-Vol-I/Web-Services-Reuse.html> 28 March 2008.
30. Poulin Jeffery (2006). "The ROI of SOA Relative to Traditional Component Reuse," Logic Library.
31. Mili, H., A. Mili, S. Yacoub and Addy, E., (2002). Reuse-Based Software Engineering: Techniques, Organization and Controls. John Wiley and Sons Ltd.
32. Linthicum, D. (2007). "How Much Will Your SOA Cost?," SOAInstitute.org, Mar. 2007. [Online]. Available: <http://www.soainstitute.org/articles/article/article/how-much-willyour-soa-cost.html>. [Accessed: Nov. 2011].
33. Knuth, D. E. (1998). The Art of Computer Programming: Volume 3, Sorting and Searching, 2nd ed.. Reading, MA: Addison-Wesley Professional May 1998.
34. Bai, Y. and Ward, R. C. (2007). "A Parallel Symmetric Block-Tridiagonal Divide-and-Conquer Algorithm," Transactions on Mathematical Software (TOMS), vol. 33, Aug. 2007, pp. A25, doi: 10.1145/1268776.1268780.
35. Khalilian, M. F., Boroujeni, Z., Mustapha, N., and Sulaiman, M. N. (2009). "KMeans Divide and Conquer Clustering," Proc. the 2nd International Conference on Computer and Automation Engineering (ICCAE 2009), IEEE Press, Mar. 2009, pp. 306-309, doi: 10.1109/ICCAE.2009.59.
36. Lin, T. Y. (2009). "Divide and Conquer in Granular Computing Topological Partitions," Proc. Annual Meeting of the North American Fuzzy Information

- Processing Society (NAFIPS 2009), IEEE Press, Jun. 2005, pp. 282-285, doi: 10.1109/NAFIPS.2005.1548548.
37. Hu F., and Wang, G. (2008). "Huge Data Mining Based on Rough Set Theory and Granular Computing," Proc. IEEE/WIC/ACM International Conference on Web Intelligence and Intelligent Agent Technology (WI-IAT'08), IEEE Press, Dec. 2008, pp. 655-658, doi: 10.1109/WIIAT.2008.84.
38. Zhang, C. and Xue, B. (2009). "Divide-and-Conquer: A Bubble Replacement for Low Level Caches," Proc. the 23rd International Conference on Supercomputing (ICS'09), ACM, Jun. 2009, pp. 80-89, doi: 10.1145/1542275.1542291.
39. Norfolk, D. (2007). "SOA Innovation and Metrics," IT-Director.com, Dec. 2007. [Online]. Available: <http://www.itdirector.com/business/change/content.php?cid=10146>. [Accessed: Nov. 2009].
40. Chaos, D. (2009). "SOA is not dead, but complexity is killing some implementations," Technoracle (a.k.a. "Duane's World"), Jan. 2009. [Online]. Available: <http://technoracle.blogspot.com/2009/01/soa-is-not-dead-but-complexity-is.html>. [Accessed: Jul. 2009].
41. Krafzig, D., Banke, K., and Slama, D. (2004). Enterprise SOA: Service-Oriented Architecture Best Practices, Upper Saddle River: Prentice Hall PTR, Nov. 2004.
42. Davies, J., Schorow, D., Ray, S. and Rieber, D. (2008). The Definitive Guide to SOA: Oracle Service Bus, 2nd ed.. New York: Apress, Sept. 2008.
43. Yusuf Lateef Oladimeji, Olusegun Folorunso, Akinwale Adio Taofeek and Adejumobi, I.A. (2011). "Visualizing and Assessing a Compositional Approach to Service-Oriented Business Process Design Using Unified Modelling Language (UML)". Computer and Information Science Journal, Vol. 4, No. 3; May 2011.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Novel Approach for Always Best Connected in Future Wireless Networks

By Kailash Chander, Dr. Dimple Juneja

Maharishi Markendeshwar University, Mullana, Haryana, India

Abstract - Basically, Vertical handover (VHO) decision relies on the selection of the 'best' available network that could meet the QoS requirements for the end-user. Therefore, a network selection mechanism is required to help mobile users choose the best network; that is, one that provides always best connected (ABC) that suits users needs and is able to change dynamically with the change in conditions. The definition of best depends on a number of different aspects such as user personal preferences, device size and capabilities, application requirements, security, present network traffic, and network signal strength. This work proposes to assign weight to all the above stated aspects so as to compute ABC. The novelty of this work is to exploit intelligent agents for weight calculations after analyzing the explored parameters for various networks. An analysis and a comparison of both services and factors for different networks are also provided in the paper.

Keywords : ABC networks, Weights and Rewards, Intelligent Agents, Quality of Service (QoS), Vertical Handover).

GJCST Classification : C.2.1



Strictly as per the compliance and regulations of:



A Novel Approach for Always Best Connected in Future Wireless Networks

Kailash Chander^α, Dr. Dimple Juneja^α

Abstract - Basically, Vertical handover (VHO) decision relies on the selection of the 'best' available network that could meet the QoS requirements for the end-user. Therefore, a network selection mechanism is required to help mobile users choose the best network; that is, one that provides always best connected (ABC) that suits users needs and is able to change dynamically with the change in conditions. The definition of best depends on a number of different aspects such as user personal preferences, device size and capabilities, application requirements, security, present network traffic, and network signal strength. This work proposes to assign weight to all the above stated aspects so as to compute ABC. The novelty of this work is to exploit intelligent agents for weight calculations after analyzing the explored parameters for various networks. An analysis and a comparison of both services and factors for different networks are also provided in the paper.

Keywords : ABC networks, Weights and Rewards, Intelligent Agents, Quality of Service (QoS), Vertical Handover).

I. INTRODUCTION

The ultimate goal of ABC is to provide QoS to the end users where QoS is defined with respect to various parameters as given in Table 1.

Table 1 : QoS Parameters

Parameters	Meaning
Speed	The time period required for a packet to reach its destination
Reliability	Reliability depends on the Bit Error Rate (BER)
Packet Loss	The probability of loss of packets during transmission.
Signal Strength	Available signal strength during transaction.
Services	Type of services supported by a network.

In addition to the above mentioned parameters A network selection also depends upon the type of service (such as Internet Surfing, Voice Data, and Video streaming) it offers, which in turn depends upon various factors such as Cost of Service, Data Rate, Mobility of Mobile Node, Signal Strength, Present Network Traffic, Security Parameter, and Drainage rate of Battery.

Author ^α : Research Scholar, Maharishi Markendeshwar University, Mullana, Haryana, India. E-mail : kailashchd@gmail.com

Author ^α : Professor, M.M Institute of Computer Technology and Business Management (MCA) , Maharishi Markendeshwar University, Mullana, Haryana, India. E-mail : dimplejunejagupta@gmail.com.

It is obvious that selection of network is initially dependent on the end user requirements, but since a user is provided with many choices, the choice becomes QoS dependent. An obvious choice would be the network that offers maximum QoS. This work analyzes QoS parameters for the type of service offered by network and factors on which a network selection depends. The work aims to assign weight to each parameter and factor as well. Our earlier work proposed deployment of agent in 4G networks and hence we propose that deployed agents shall do all weight computations, reducing the overhead of service providers and enhancing the QoS to service users.

The paper has been organized in four sections. Section II presents the related work. Section III discusses the proposed solution and compares the results. Finally conclusions and future scope are presented in Section IV.

II. RELATED WORK

In fourth generation (4G) communication [26], selection of the best access network is the major challenge for future research. Many researchers [19, 20, 21, 22] proposed different solution for achieving "Always Best Connected". Authors [5] proposed Game theoretical modelling to solve the network selection problem. Game theory is the study of a mathematical model of conflict and co-operation between intelligent reasonable decision makers [6]. According to a Game theoretical Model [7], the players of the game are the individual access network (WLAN, GSM, WCDMA, WiMAX, WiFi etc.) each of which contends to win a service request.

Another solution to select an appropriate network is based on distance function [8]. Distance function generates an ordered list of various access technologies called networks in a particular region according to multiple user preferences and level of interest. The proposed algorithm works on the user-specified parameters i.e. Bandwidth Utilization, Call Drop during Handoff, Cost of Services, Battery Power etc.

Work in [9] provided an efficient load balancing based access point selection algorithm which consider the direction of advancement of the mobile node and hence is able to extract the best possible network for the user equipment to link up, as it moves.

Exploiting ants in telecommunication [23, 24, and 25] has been the interest of researchers and it has

been proved that such agent based frameworks have been instrumental in improving the performance of existing networks.

Weighted distance function [10] is obtained based on multiple QoS parameters as per user needs. The proposed algorithm shows better results compared to single parameter based system, under a heterogeneous network system.

Another solution of network selection is through QoS Broker [11]. This Broker monitors the QoS performance actively for each wireless network, and then the result of this monitoring will be passed to analysis statistics of all the QoS parameters in each network to get the best network.

In [12] authors developed a process to evaluate three packet-switched networks (UMTS, WLAN and GPRS) in reference to the QoS offered. It also identifies the weak points of a network and finally selects the network that offers the highest standard for QoS.

Further, focusing on the use of mobile agents in telecommunication section, Literature [1, 2, 15] indicates communication applications are modeled as a collection of agents and each agent occupies different locations at different times since it can move from one place to another.

Mobile AGeNt Architecture (MAGNA) [2] is being developed by GMD-Fokus for future telecommunications applications, in which the conventional client server concept is cordially complemented by agent concepts. This framework is used for the development of agent-based telecommunications applications which exhibit rapid, decentralized provisioning of intelligent services on demand.

Among the different paradigms of intelligent agents, Reinforcement Learning (RL) [17] appears to be particularly appropriate to address a number of the challenges of the future mobile communication. RL involves learning what to do and how to map situations with actions to maximize a numerical value signal [3]. A Reinforcement Learner Agent (RLA) ascertains on its own which actions to take to get the maximum weight value. The agent learns from its mistakes and come up with a policy based on its experience to maximize the attained weight value [4].

Focusing our attention to agent-based solutions, it is apparent that there is a scope of an agent based solution for ABC network in future. Next section presents such a solution.

III. PROPOSED SOLUTION

The process of network selection refers to the process of deciding over which network to connect at any point in time. On the other hand user wants the selection of service among the available networks according to his/her requirements. Thus a novel network

selection mechanism is being proposed such that the selected network satisfies the current session's QoS requirements. Each network would be assigned a weight which is based on the QoS parameters and factors that it provides and satisfies the end user requirements. The agent is then required to compute the sum of all weight assigned to a particular network which is then normalized within the range of 0-1. A network scoring maximum (i.e. 1) shall be the best available network while a network scoring less than a specified threshold ($<.5$) shall be ignored and similarly, a network gaining a zero weight shall be straightaway discarded. Following section presents the rule set that must be followed by agent for assigning weights.

Rule Set

- a. **Cost of Service** : Maximum value is given if cost of service is either zero or very less, because every user wants better service always at lowest price.
- b. **Data Transfer Rate** : Present data transfer speed of the available network.
- c. **Mobility of Mobile Node**: The value will be calculated on the basis of present status of movement of mobile node. If mobile is static the maximum value i.e. 1.0 is assigned and if node is in moving state then 0.0 is allocated.
- d. **Signal Strength** : It involves the minimum or maximum signal strength supported by a particular network. It will provide us the net signal strength of that network.
- e. **Present Network Traffic** : Every network provides services to its user based on contention ratio such as 1:1 or 1:10 or 1:50. Therefore the total bandwidth is divided based on this ratio. If a network is providing service with a content ratio of 1:20 and presently only four users are logged in then user will get more bandwidth and higher speed.
- f. **Security Parameter** : A secure network is always treated as a best network because of security threats. Selecting best secure network is a challenging task. A good network is that which supports maximum security layers. Some of security layers are Network Intrusion Detection System, Firewall, Email Scanning, Internet Security, Server Level Virus Scanning, Workstation Virus Scanning, and Updated Communication Software. Agents calculate the weight value based on number of layers supported by a network.

Table 2 : Weight Rules

Parameters Weight (W) ↓	Cost of Service Offered (A)	Data Transfer Rate (DTR) (B)	Mobility of Node (C)	Signal Strength (D)	Network Traffic (E)	Security (F)	Drainage Rate of Battery (G)
1.0	Zero	11mbps-100mbps	Static	Excellent	Very few users/Single User	Fully Secure	Very Light Application
0.75	Negligible	2mbps-11mbps	Walking	High	Moderate	High	Light Application
0.50	Moderate	128kbps-2mbps	--	Good	High	Moderate	Moderate
0.25	High	<128kbps	--	Low	Very High	Low	Heavy
0.0	Unaffordable	--	Very High Mobility	Nil	Extreme (All routes are busy)	None	Very Heavy Application

g. **Drainage rate of Battery** : An algorithm in [15] proposes to shift to the lesser power demanding network in case the present battery status of mobile node is not sufficient for current transaction. Here, again the value is calculated on the basis of consumption of battery life for a particular application for a particular network by comparing present battery life of mobile node i.e. drainage of battery is application dependent for instance heavy application implies more consumption of battery.

The weights assigned as per the rules mentioned above are being listed in Table 2. The above factors compute the weight-age for available networks. Some of the above factors have higher impact while others have less for network selection decision. For an ABC network, the user agent computes the sum of all weights (computed as per the rule set given above) as per the following formula:

$$ABC = \sum_{i=1}^n W_i$$

Where, n represents the numbers of parameters taken into consideration and W_i represents the corresponding weight computed for a respective parameter.

IV. SEQUENCE DIAGRAM

A user agent is invoked whenever a user demands for ABC network. An individual user agent cannot evaluate the complete rule set for different types of network. Therefore, it broadcasts the request and in response to this request, various network agents respond with bid to the user agents.

The bid comprises of weights which have been computed on the basis of rule set as already defined. User agent then computes the sum of weights to find out the ABC network shown in Figure 1. It is obvious that a network having maximum weight will be selected and

rest would be discarded. Each network agent is responsible for providing the data on the basis of their rule set and all agents are set to move and execute state as every network service provider would intend to provide service to a mobile user.

The following algorithm for MAGagent given below provides the details of working of the proposed framework for its implementation.

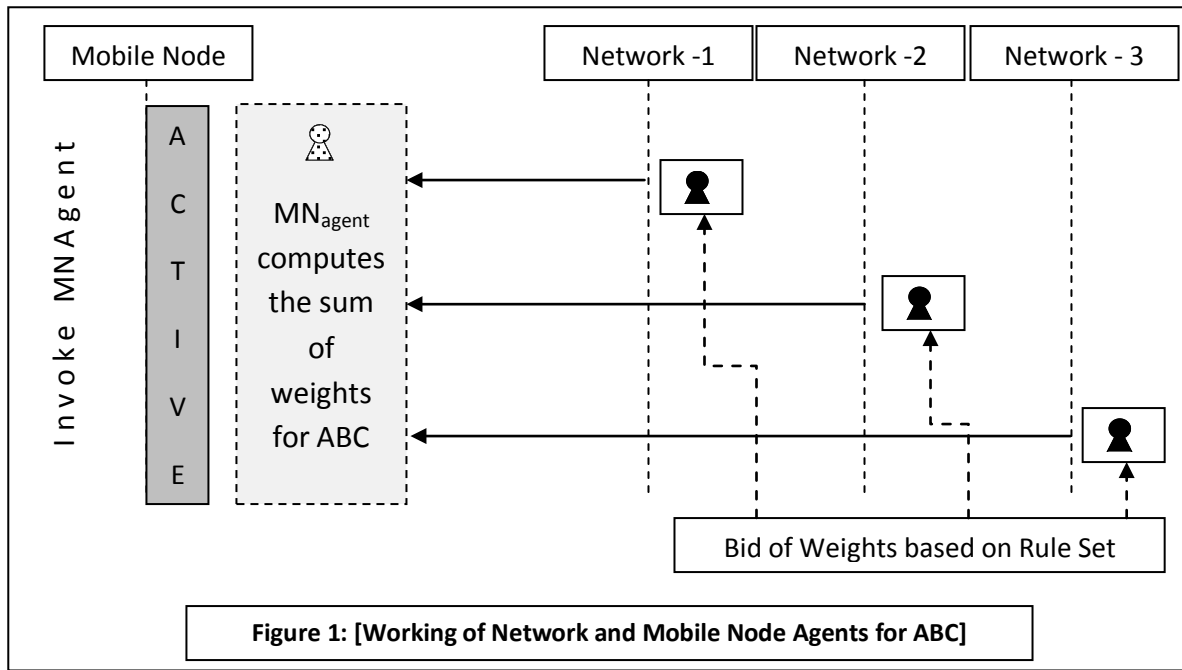
Algorithm :

Algorithm: MAG_{agent}

```

1: begin
2:   i=0
3:   invoke MNagent
4:   aN[ ] = <Available Networks>
5:   while (i <= aN[ ] )
6:     j=0
7:     while (j <= 7)
8:       Invoke MAGagent[j]
9:       ds= getDataSet(f)
10:      wi= <assign as per Rule
Set>>
11:      sw[j]=sw[j]+wi
12:      j++;
13:     wend
14:     MNagent ← aN[i].sw[j];
15:     i++;
16:   wend
17:   ABC= sort (aN[ ].sw[ ]) → descending order
18: end

```



V. DEMONSTRATION AND PROOF OF CONCEPT

In this section the concept is demonstrated by taking different possible weight values for the parameters and sum of all the parameter values is computed to decide on the best network at any point. The algorithm will process the input provided in this format and will decide which network can provide best services to the subscriber at that time.

CASE-1								
(Available Network)	Parameters							
	(A)	(B)	(C)	(D)	(E)	(F)	(G)	(Sum)
Net – 1	0.75	0.50	1.00	0.50	0.50	0.25	0.50	4.00
Net – 2	1.00	0.50	1.00	0.25	0.25	0.50	1.00	4.50
Net – 3	0.25	1.00	1.00	1.00	0.75	0.75	0.25	5.00
CASE-2								
(Available Network)	Parameters							
	(A)	(B)	(C)	(D)	(E)	(F)	(G)	(Sum)
Net – 1	0.25	0.75	0.75	0.50	1.00	1.00	0.5	4.75
Net – 2	0.75	0.50	0.75	0.50	0.75	0.75	0.5	4.50
Net – 3	0.50	0.75	0.75	0.50	0.75	0.50	0.25	4.00

In the first case the mobile node is in static mode, thus Net-3 will be selected since it is providing high data at low cost while in the second case mobile node will prefer Net-1 because of its security although the cost of network is high. In the second case status being mobile the node will prefer secure transaction, even with more cost.

VI. CONCLUSION

In future 4G mobile environments, various access technologies will coexist, complementing each other. Therefore, a network selection mechanism is required to help mobile users choose the best network; that is, one that provides always best connected (ABC) that suits users needs, and changes, if conditions changes. Thus a novel network selection mechanism using intelligent agents has been proposed, which select the best network based on QoS parameters. The security aspect of agents has been ignored and would be taken up in future works.

REFERENCES REFERENCES REFERENCIAS

1. Marie-Pierre Gervais and et.al. "Enhancing Telecommunications Service Engineering with Mobile Agent Technology and Formal Methods" IEEE Communications Magazine, July 1998.
2. S. Krause and T. Magedanz, Mobile Service Agents Enabling, "Intelligence on Demand in Telecommunications," Proc, IEEE GLOBECOM '96, London, UK, Nov. 1996.
3. Abdulhai, B. and Kattan L. "Reinforcement learning: Introduction to theory and potential for transport applications" Canadian Journal of Civil Engineering, Volume 30, Number 6 (11), 981-991. 2003.
4. Sutton, Richard S. and Barto, Andrew G. . Reinforcement Learning. MIT Press, Massachusetts USA, 322 pp. 2000.
5. Mohammad Sazid Zaman Khan et. al. "A Network Selection Mechanism for Fourth Generation Communication Networks" Journal of Advances in Information Technology, Vol.1 No.4 p.p. 189-196, 2010.

6. Roger Myerson "Game Theory: Analysis of Conflict," Harvard University Press, 1997.
7. J. Antoniou, A. Pitsillides, "4G Converged Environment: Modeling Network Selection as a Game," Proc. of IST Mobile Summit 2007, Budapest, Hungary.
8. T. Nisha Devi et. al. "A Novel Approach for Optimal Network Selection in Heterogeneous Wireless Networks", Proceedings of the International Joint Journal Conference on Engineering and Technology (IJJCET) p.p. 158-161, 2010.
9. I.S. Misra, A. Banerjee "A Novel Load Sensitive Algorithm for AP selection in 4G Networks"
10. Amit Sehgal, Rajeev Agrawal, "QoS Based Network Selection Scheme for 4G Systems", IEEE Transactions on Consumer Electronics Vol. 56, No. 2, p.p. 560-565, May 2010.
11. Mr. Mohammed Saeed Jawad, et. al. "Optimizing Network Selection to Support End-User QoS Requirements for Next Generation Network" International Journal of Computer Science and Network Security, Vol. 8 No.6, p.p. 113-117, June 2008.
12. Dimitris Charilas et. al. "Packet-switched network selection with the highest Qos in 4G networks" Computer Networks 52(2008) p.p. 248-258, Science Direct, Elsevier.
13. Chen Yiping and Yang Yuhang "A New 4G Architecture providing Multimode Terminals Always Best Connected Services" IEEE Wireless Communications, 2007 p.p. No 36-4.
14. EVA Gustafsson and Annika Jonsson, "Always Best Connected" IEEE Wireless Communications, 2003 p.p. no. 49-55.
15. Kailash Chander, Dr. Dimple Juneja "Mobile Agent based Emigration Framework for 4G: MAEF " communicated to International Journal of Information and Computing Science, UK.
16. Torsten Braun "Ant-Based Mobile Routing Architecture in Large-Scale Mobile Ad-Hoc Networks".
17. Richard S. Sutton "Generalization in Reinforcement Learning: Successful Examples Using Sparse Coarse Coding.
18. F. Ahdi, B. H. Khalaj, "Vertical Handoff Initiation Using Road Topology and Mobility Prediction". IEEE 2006.
19. Gabor Fodor, Anders Eriksson, et. al. "Providing Quality of Service in Always Best Connected Network". IEEE July 2003.
20. Chen Yiping And Yang Yuhang, "A New 4G Architecture Providing Multimode Terminals Always Best Connected Services". IEEE Wireless Communications, April 2007.
21. Tein-Yaw Chung, Fong-Ching Yuan, "Extending Always Best Connected Paradigm for Voice Communications in Next Generation Wireless Network". IEEE 2008.
22. Yiping Chen, Chen Deng, "Access Discovery in Always Best Connected Networks". IEEE 2008.
23. Lu guoying Zhang subing Liu zemin "Distributed Dynamic Routing Using Ant Algorithm For Telecommunication Networks". IEEE 2000.
24. Mohammad Mursalin Akon and Dhrubajyoti Goswami, "A New Routing Table Update and Ant Migration Scheme for Ant Based Control in Telecommunication Networks". Proceedings of the 7th International Symposium on Parallel Architectures, Algorithms and Networks (ISPAN'04), IEEE 2004.
25. Heesang Lee, Gyuwoong Choi, "An Ant-Assisted Path-Flow Routing Algorithm for Telecommunication Networks". IEEE 2004.
26. Ronal van Eijk, Jacco Brok, Jeroen van Bommel and Bryan Busropan, "Access Network Selection in a 4G Environment and the Roles of Terminal and Service Platform" In Wireless World Research Forum.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Improved Energy and Latency Efficient MAC Scheme for Dense Wireless Sensor Networks

By K. P. Sampooram, Dr. K. Rameshwaran

K.S.R. College of Engg., Tiruchengode

Abstract - The sensor nodes in the wireless sensor networks have limited battery power, which motivates to work on energy conserved MAC schemes for better lifetime and latency efficient. Previous work carried out in energy conserved MAC schemes are limit the idle listening time, reduces overhearing (sensor node hear a packet destined for other nodes) and minimizing the used control packet size. The current existing work presented ELE-MAC (i.e. Energy Latency Efficient MAC) which adopts less control packets to preserve energy in sparsely distributed sensor nodes of the wireless sensor networks. It performs statistically the same or better latency characteristic compared to adaptive SMAC. ELE-MAC follows the adaptive listening technique, which reduce the sleep delay introduced by the periodic sleep of each node in case of a multi-hops network. The proposal in this work, extends the ELE-MAC to work efficiently with wireless sensor network comprises of high node density by combining the RTS and SYNC control packets. The extended version uses two separate frequencies for data and control packets to avoid the use of handshake mechanisms (e.g. RTS/CTS) in order to reduce energy consumption and packet delay. It enables a receiver to send a busy tone signal on the control channel and notify the neighbors about the ongoing reception of data in progress. This process avoids packet collisions and in turn improves the node lifetime and throughput. The nodes in a sensor network have their own different traffic loads according to the tasks assigned and their locations. The extension of ELE MAC adopts the different traffic loads of each node as performance metric for reducing the latency. Each sensor node calculates its utilization after the last synchronization time, and adjusts its duty cycle according to the calculated utilization, and then send new schedule to its neighbors via broadcasting.

Keywords : *Wireless Sensor Network, Lifetime, Traffic load, Energy and Latency, MAC.*

GJCST Classification : *C.2.1*



Strictly as per the compliance and regulations of:



Improved Energy and Latency Efficient MAC Scheme for Dense Wireless Sensor Networks

K. P. Sampoornam^a, Dr. K. Rameshwaran^a

Abstract - The sensor nodes in the wireless sensor networks have limited battery power, which motivates to work on energy conserved MAC schemes for better lifetime and latency efficient. Previous work carried out in energy conserved MAC schemes are limit the idle listening time, reduces overhearing (sensor node hear a packet destined for other nodes) and minimizing the used control packet size. The current existing work presented ELE-MAC (i.e. Energy Latency Efficient MAC) which adopts less control packets to preserve energy in sparsely distributed sensor nodes of the wireless sensor networks. It performs statistically the same or better latency characteristic compared to adaptive SMAC. ELE-MAC follows the adaptive listening technique, which reduce the sleep delay introduced by the periodic sleep of each node in case of a multi-hops network.

The proposal in this work, extends the ELE-MAC to work efficiently with wireless sensor network comprises of high node density by combining the RTS and SYNC control packets. The extended version uses two separate frequencies for data and control packets to avoid the use of handshake mechanisms (e.g. RTS/CTS) in order to reduce energy consumption and packet delay. It enables a receiver to send a busy tone signal on the control channel and notify the neighbors about the ongoing reception of data in progress. This process avoids packet collisions and in turn improves the node lifetime and throughput.

The nodes in a sensor network have their own different traffic loads according to the tasks assigned and their locations. The extension of ELE MAC adopts the different traffic loads of each node as performance metric for reducing the latency. Each sensor node calculates its utilization after the last synchronization time, and adjusts its duty cycle according to the calculated utilization, and then send new schedule to its neighbors via broadcasting.

Keywords : Wireless Sensor Network, Lifetime, Traffic load, Energy and Latency, MAC

I. INTRODUCTION

Wireless sensor network (WSN) is a collection of sensor nodes that interact with each other intentionally to gather information from the surveillance area. Sensor nodes support unattended operation for long duration, usually in remote areas. WSN applications such as environmental monitoring [1], object tracking [2] and intelligent buildings [3] require a reliable data transmission and can endure long periods of operation. The limitation of sensor nodes are low

processing capabilities (delay attenuation), low power battery and low memory capacities which initiates to improve those constraints in increasing the network life time of wireless sensor networks.

Medium access control (MAC) layer manages the medium accessibility to minimize collision among transmitting packets. Packet collision requires node to retransmit the packet, hence consuming additional energy. MAC layer controls the physical (radio transceiver) layer which has greater effect on overall energy consumption and lifetime of a node. The nodes sometimes falsely assumed that the channel is in idle condition and start the transmission which results in data collision lead to more energy requirement. In some cases nodes are exposed due to out of receiver range which leads to overhearing and increases the delay of transmission. In few other cases nodes receive one of two simultaneous transmissions, which creates complex traffic load control. The idle listening of a sensor node due to continuous listening of the channel to receive a potential packet from its neighboring nodes consumes more energy.

Energy inefficiency caused by the idle-listening problem and high collision probability can be avoided in Time Division Multiple Access (TDMA) based MAC protocols. The existing work presented an ELE-MAC Energy Latency Efficient MAC with distributed TDMA mechanism which possesses an active/sleep mechanism for efficient energy usage with predefined duty cycle. However the ELE-MAC work on different traffic load condition affects the overall network lifetime of the sensor network. In this paper, we presented an improved version of ELE-MAC which works on balancing the different traffic load conditions of the sensor node transmission on the target object being detected.

II. LITERATURE REVIEW

MAC protocol is classified into random access and conflict-free multiple access. Traditional MAC protocols such as ALOHA [5], CSMA [6], and MACA [7], are designed based on contention based random access approach. The classic ALOHA protocol uses simple transmission mechanism where node transmits a packet when it is generated. However, its simplicity comes at an expense of very high probability of packet collision; hence increases the energy expenditure due to packet retransmission. Therefore, Carrier Sense Multiple

^a Author : Department of ECE, K.S.R. College of Engg., Tiruchengode-637215 Telephone : 9629377780 E-mail : sambu_swasti@yahoo.com

^a Author : Principal, JJ College of Engg, Tiruchirappalli – 620 009. Telephone : 9952266566 E-mail : krameshwaran@gmail.com

Access (CSMA) protocol is developed [6] with the objective of minimizing collision by implementing a small time for channel listening in order to detect channel activity. However, the protocol cannot solve the hidden terminal problem which normally occurs in ad-hoc networks where the radio range is not large enough to allow communication between arbitrary nodes and two or more nodes may share a common neighbor while being out of each other's reach. The MACA protocol introduces a three-way handshake mechanism to make hidden nodes aware of upcoming transmission, so collision at neighboring nodes can be avoided. However, the handshaking mechanism causes overhead on control packet.

All these protocols require all nodes to continuously listen to the channel due to unpredictable packet transmission by its neighboring nodes, hence introducing a problem called idle-listening problem. This situation causes a node to expend a lot of wasteful energy causing the implementation of these protocols in WSN inefficient. Sensor-MAC (SMAC) protocol [8] attempted to solve the problems by introducing active-sleep cycles in the presence of random access channel. Node will execute a variant of MACA contention-based MAC protocol during active period to minimize the hidden terminal problem, while turning its radio off during sleep period to reduce idle listening problem. However SMAC implements neighbors' information variables called Network Allocation Vector (NAV) [9] for its collision avoidance technique. Node checks the NAV value before sending the RTS message. Nevertheless, implementing contention based mechanism is still vulnerable to collision due to random mechanism in its data packet transmission.

Energy inefficiency caused by the idle-listening problem and high collision probability can be avoided in Time Division Multiple Access (TDMA) based protocols. In TDMA-based protocol such as HiperLan-II [10], time is divided into several frames, and a frame is divided into a number of time slots. Since all transmissions within the frame are pre-scheduled, it is possible for a node to sleep when it is not expected to transmit or receive any packets. Thus, the TDMA-based MAC protocol can clearly avoid the over-emitting problem. Since only the owner of the time slot is allowed to transmit a packet, collision problem can be avoided significantly.

Tahar et al.,[1] presented an energy efficient MAC protocol which realizes both energy efficiency and improve the channel utilization compared to the already existed techniques. For this they provided ELE-MAC with the inputs from adaptive SMAC scheme. In this a control packet strategy which presented a packet exchange sequence aiming to minimize the energy wasted by control packets and to decrease latency.

However different traffic load condition of the sensor nodes on transmission of target detected objects at any given time also further increases the latency and energy. The proposal in this work present an another variant of ELE-MAC which handles the different load conditions based on load distribution and scheduling mechanism of each and every sensor nodes of the network to improve the overall lifetime.

III. LOAD BALANCED ELE-MAC

In the wireless sensor network, the control packet has greater impact on the network power consumption which is comparable to the size of data packets. Energy consumption is reduced by optimizing the exchanged control packets. This motivates to present energy efficient MAC protocol that minimizes the exchanged control packets (ELE-MAC) compared to that of adaptive SMAC protocol. ELE-MAC control packets shown in Fig 2 as compared with normal control packets (in Fig 1) present the operations of how energy conservation happens in wireless sensor networks.

The ELE-MAC control packet provides two additional fields (i.e. ACK destination Node Address and ACK flag) which allow the new RTS packet to play the role at the same time of an ACK and a RTS. This new packet will be exchanged only when data are sent adaptively (i.e. not at the scheduled listen time). Thus, no ACK packet will be emitted in that case. The transmission is performed normally (i.e. at the scheduled listen time). Each data packet received is followed by an ACK to the sender.

The operation of ELE-MAC shown in Fig 3 explains the operation that node A has data to be transmitted to node B to end in node C which is the sink of the illustrated topology. The ELE-MAC scheme starts the adaptive wake up period immediately after receiving the data packet instead of waiting for the ACK packet like for the SMAC adaptive listening mechanism. This modification is made for allowing a receiver to inform its neighbors about the data reception through the ACK flag field. Also, this packet allows the receiver to mention its need to transmit the received data packet to the next-hop if it exists (i.e. send RTS).

The most common workload in sensor networks consists on small periodic data packets. Thus, ELE-MAC doesn't propose a fragmentation mechanism. Like IEEE 802.11 and SMAC, broadcast packets are sent only when virtual and physical carrier sense indicate that the medium is free. In addition, these packets will not be preceded by RTS/CTS and will not be acknowledged by their recipients.

The load balance scheme proposed for ELE-MAC to multi-hop multi-channel sensor networks. Based on the load distribution of all sensor nodes, it algorithm

dynamically alternates the communication channels. As a result, the extra load from over-loaded channels is directed to under-loaded channels with a computed switch probability.

In addition a high throughput is achieved with stabilized load conditions on the sensor nodes during the transmission of more number of target objects being detected. The performance of the load balance algorithm is evaluated through simulation studies on both ELE-MAC of energy variant and load variant.

Original RTS Packet

Type	Length	Destination node address	Source node address	Duration	CRC
------	--------	--------------------------	---------------------	----------	-----

Fig 1: Wireless Sensor Network Normal RTS Control Packet

ELE-MAC RTS Packet

Type	Length	RTS destination node address	ACK destination node address	Source node address	ACK Flag	Duration	CRC
------	--------	------------------------------	------------------------------	---------------------	----------	----------	-----

Fig 2 : Wireless Sensor Network ELE-MAC RTS Control Packets

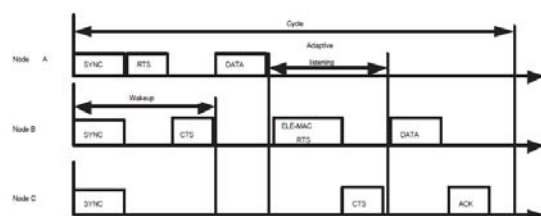


Fig 3 : ELE-MAC Operations

Type	Length	RTS Destination node address	Target object load rate	ACK Destination node address	Source node address	ACK flag	Load flag	Duration	CRC
------	--------	------------------------------	-------------------------	------------------------------	---------------------	----------	-----------	----------	-----

Fig 5 : Load Distribution & Scheduling Packet

The routing protocol makes greedy forwarding decisions using information about a router's immediate neighbors in the network topology. In fact, to let each node hear only its next neighbor, we put nodes distant by 100 meters taking into account that the transmission range in NS-2 is set to 150 meters. No mobility is assumed in our simulation scenarios. As the goal of our simulation is to compare the performance of SMAC with ELE-MAC, we choose our traffic source to be constant bit rate (CBR) source. The NS-2 ELE-MAC and load variant simulation parameters are analyzed to extract the useful traces and to compute the energy consumption as well as the latency with TCL scripts. The simulation is



Fig 4 : Adaptive SMAC

IV. EXPERIMENTAL EVALUATION OF IMPROVED ELE-MAC

The simulation is carried for improve load balanced ELE-MAC with existing ELE-MAC in NS-2. By comparing with SMAC (in its two alternatives), it is shown that load balance variant of ELE-MAC shows better lifetime of the sensor network in terms of load stability, latency and energy consumption. The adaptive SMAC implementation deployed in this NS's version doesn't provide us with the correctly nodes' energy consumption. Further, the problem resides in the implemented Energy Model. This is because it doesn't take into consideration the energy wasted by idle listening (i.e. doesn't drain energy in the sleep/wakeup methods). Henceforth, to enable the right tracking of the energy consumed by each node at any time, we tune the energy model and the SMAC sources.

The behavior of the proposed load variant ELE-MAC when varying the traffic load and because of the limited transmission range of wireless network interfaces (i.e. multiple network hops may be required for one node to exchange data) a multi-hops environment is required. Similar to the test bed realized for evaluating SMAC on a multi-hop networks, a linear topology composed from 40 to 60 nodes with only one source and a sink which is chosen the later node in the multi-hops chain. This simple topology allows us to concentrate on the inherent properties of load variant and energy variant ELE-MAC.

carried out for several pause time to obtain significant statistical results.

The Control packets analysis show the resultant of energy variant and load variant ELE-MAC with different traffic rate sources on the wireless sensor network. Energy variant ELE-MAC exchanges few control packets compared with load variant initially however on continuous simulation load variant shows better result than energy variant. Fig.6 presents the energy consumption performance of Energy variant and Load variant ELE-MAC. Conserve the energy which would be lost by the control packets overhead by maintaining the stability of load on all the sensor nodes.

The Latency analysis handles the end-to-end delay quantification from the simulation viewpoint. ELE-MAC energy and load variant measure the total time required to transmit the generated data packets. Load ELE-MAC achieves better latency performance compared to that of energy ELE-MAC.

Fig.7 plots the latency performance of Energy variant and Load variant ELE-MAC. The load distribution and scheduling policy of load variant ELE-MAC reduces the listening time of the sensor node in the MAC layer which in turn reduces overhearing and latency. As can be seen in Fig.8 the scheduling policy of load variant ELE-MAC helps the node transmission stability to its optimal level.

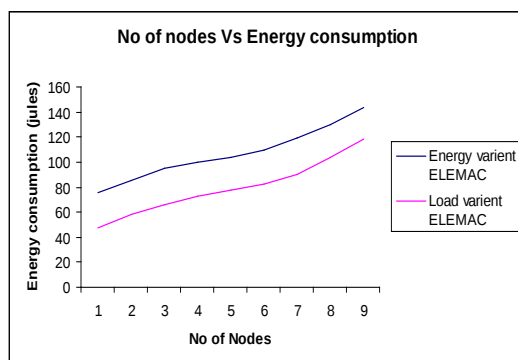


Fig 6 : Energy consumption performance

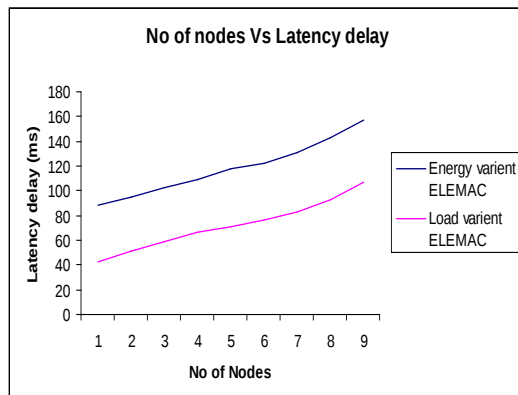


Fig 7 : Energy Latency performance

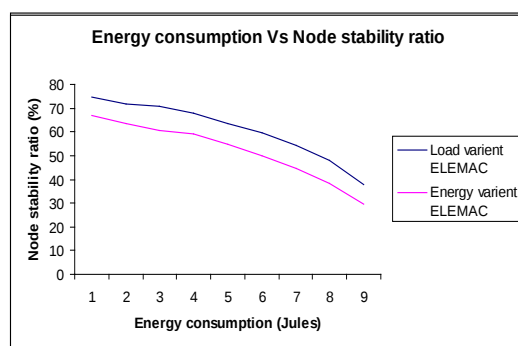


Fig 8 : Node Stability performance

V. CONCLUSION

The improved variant of ELE-MAC which balances the different load condition of the sensor nodes by load scheduling and distribution scheme improves the overall sensor network lifetime. The simulation of Load balanced ELE-MAC in the test bed consisting of 100 sensor nodes is carried to evaluate the performance of energy consumption levels, latency and load rate of each sensors. The performance results shows that the proposed load balanced ELE-MAC shows better sensor node stability (nearly 10%) in transmitting the detected target objects with less energy consumption (decreases to 15%) and latency (decreased to 23%) compared to that of existing ELE-MAC.

REFERENCES REFERENCES REFERENCIAS

1. Tahar Ezzedine, Mohamed Miladi and Ridha Bouallegue, "An Energy-Latency-Efficient MAC Protocol for Wireless Sensor Networks", International Journal of Electrical and Computer Engineering 4:13 2009.
2. Mainwaring, J. Polastre, R. Szewczyk, D. Culler, and J. Anderson. Wireless Sensor Networks for Habitat Monitoring. International Workshop on Wireless Sensor Networks and Applications. 2002, pp. 88-97.
3. W. Ye, J. Heidemann, and D. Estrin. Medium Access Control with Coordinated Adaptive Sleeping for Wireless Sensor Networks. IEEE/ACM Trans. Net., vol. 12, no. 3, June 2004, pp. 493 – 506.
4. L. V. Hoesel, P. Havinga. Collision-free Time Slot Reuse in Multihop Wireless Sensor Networks. Proc. Int. Conf. on Intelligent Sensor, Sensor Networks and Information Processing, Dec. 2005, pp. 101 – 107.
5. L. F. W. van Hoesel and P. J. M. Havinga. Design Aspects of An Energy-Efficient .Lightweight Medium Access Control Protocol for Wireless Sensor Networks. July 17, 2006.
6. J. Hill, R. Szewczyk, A. Woo, S. Hollar, D. E. Culler and K. Pister. System Architecture Directions for Networked Sensors. In Architectural Support for Programming Languages and Operating Systems. 2000. pp. 93 104.
7. D. Gay, P. Levis, R. von Behren, M. Welsh, E. Brewer and D. Culler. The NesC Language: A Holistic Approach to Networked Embedded System. In Proc. ACM SIGPLAN 2003, Conf. on Programming Language Design and Implementation. June 2003. pp. 1-11.
8. IEEE Computer Society. Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specification. IEEE Std. 802.11, 1997.
9. S. Lin, J. Zhang, G. Zhou, L. Gu, T. He and J. A. Stankovic. ATPC: Adaptive Transmission Power

- Control for Wireless Sensor Networks. In Proceedings of the 4th International Conf. on Embedded Networked Sensor Systems. Colorado. 2006. pp. 223- 236.
10. Rozeha Abd Rashid, Wan Mohd Ariff Ehsan W Embong, Norsheila Fisal, "Enhanced Lightweight Medium Access Control Protocol", Radio Frequency & Microwave Int. Conf. (RFM) KL, Dec 2-4, 2009 Proceedings of the International Conference on Circuits, Systems, Signals 301 .





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Optimal High Performance Self Cascode CMOS Current Mirror

By Vivek Pant, Shweta Khurana

Kurukshetra University Kurukshetra

Abstract - In this paper the current mirror presented, having low voltage and mixed mode structure has been proposed. The performance of self cascode MOSFET current mirror is optimized with high output impedance and can operate at 1 V or below. Simulation results conform to Analog Mentor tools having Design Architect for schematics and Eldonet for SPICE simulation, with input reference current of $20\mu\text{A}$. This review paper presents a comparative performance study of self cascode current mirror with other current mirrors.

Keywords : *current mirrors, cascode current mirror, low voltage analog circuit.*

GJCST Classification : 1.2.9



Strictly as per the compliance and regulations of:



Optimal High Performance Self Cascode CMOS Current Mirror

Vivek Pant^α, Shweta Khurana^Ω

Abstract - In this paper the current mirror presented, having low voltage and mixed mode structure has been proposed. The performance of self cascode MOSFET current mirror is optimized with high output impedance and can operate at 1 V or below. Simulation results conform to Analog Mentor tools having Design Architect for schematics and Eldonet for SPICE simulation, with input reference current of 20μA. This review paper presents a comparative performance study of self cascode current mirror with other current mirrors.

Keywords : current mirrors, cascode current mirror, low voltage analog circuit.

I. INTRODUCTION

To meet the needs of present era of low power portable electronic equipment, many low voltage design techniques have been developed. This led to the analog designers to look for innovative design techniques like Self cascode CMOS Current Mirror [1-5]. In this paper, we have investigated the merits and demerits of various current mirror configurations. For this we designed the basic current mirror first then improved our results by using various configurations like cascode current mirror, Wilson current mirror and finally the current mirror based on self cascode CMOS and analyzed its results through the SPICE simulations for 0.35 micron CMOS technology.

II. BASIC MOSFET CURRENT MIRROR

The basic current mirror can also be implemented using MOSFET transistors (Fig: 1). Transistor M1 is operating in the saturation or active mode, and so is M2. In this setup, the output current I_{OUT} is directly related to I_{REF}, as discussed next. Simulation results for I_{out} vs V_{DS} curve for Basic Current Mirror is shown in fig 2. For a current mirror, neglecting channel length modulation:-

$$I_{out} = \frac{1}{2} \mu_n C_{ox} (W/L)_2 (V_{gs} - V_{th})^2 \quad (1)$$

$$I_{ref} = \frac{1}{2} \mu_n C_{ox} (W/L)_1 (V_{gs} - V_{th})^2 \quad (2)$$

When eq. 1 is divided by eq. 2, we have

$$I_{out} = I_{ref} (W/L)_2 / (W/L)_1$$

Limitations

1. As we can see from the basic current mirror circuit current gain is poor and the output current is having the channel length modulation effects. This is verified in eq. 3

Author ^α : Electronic Science department, Kurukshetra University Kurukshetra.

$$I_{out} = I_{ref} \frac{(W/L)_2 (1 + \lambda V_{ds2})}{(W/L)_1 (1 + \lambda V_{ds1})} \quad (3)$$

Here $V_{ds1} \neq V_{ds2}$.

2. Output resistance is finite and small value.

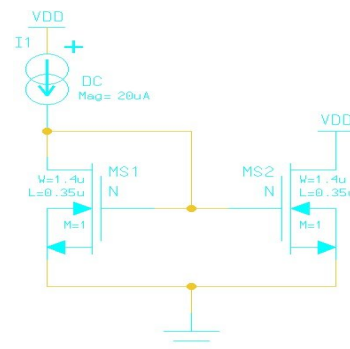


Fig 1: Basic current mirror

Simulation results:

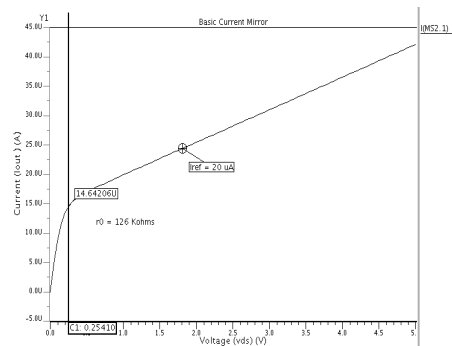


Fig 2 : I_{out} vs V_{ds} curve for Basic Current Mirror

III. CASCODE CURRENT MIRROR

The idea of cascode structure is employed to increase the output resistance (Fig.3) and the implementation requires NMOS technology. It is used to remove the drawback of channel length modulation in basic current mirror. In the simulation results of basic current mirror the channel length modulation effect was not considered. In practice, this effect results in significant error in copying currents. The circuit features a wide output voltage swing and requires an input voltage of approximately one diode drop plus a saturation voltage. By maintaining the input transistors in saturation, the output current will track the input current, regardless of increases in ambient temperature [6, 7, 8].

Simulation results for I_{out} vs V_{ds} curve for Cascode Current Mirror are shown in fig 4.

Advantages:

1. Cascode current mirror eliminates the channel length modulation effect by keeping $V_{ds1} = V_{ds2}$ constant in the ratio:

$$I_{out} = I_{ref} \frac{(W/L) (1 + \lambda V_{ds2})}{(W/L) (1 + \lambda V_{ds1})}$$

2. Improves output resistance.

Disadvantages

1. Less accurate.
2. Current becomes constant for quite large value of V_{ds} e.g. in this case minimum V_{ds} is 1.2 V.
3. Body effect is also present which disturbs the output current.

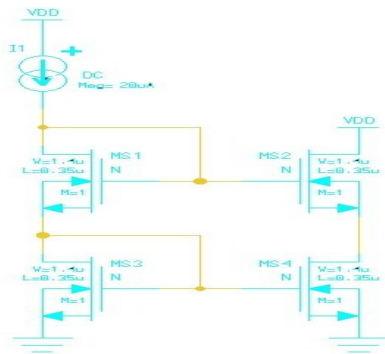


Fig 3 : cascode current mirror

Simulation results :

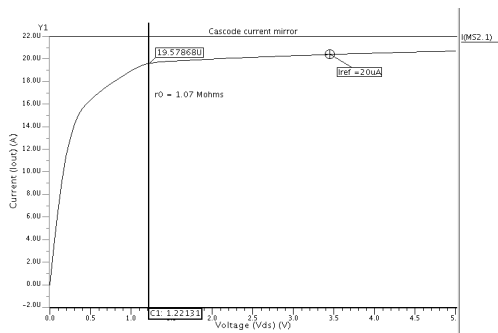


Fig 4 : I_{out} vs V_{ds} curve for Cascode Current Mirror

IV. WILSON CURRENT MIRROR

A Wilson current mirror or Wilson current source is a circuit configuration designed to provide a constant current (Fig:5). This circuit has the advantage of virtually eliminating the current mis-match of the conventional current mirror thereby ensuring that the output current I_{out} is almost equal to the reference or input current I_{Ref} thus eliminating the drawbacks of cascode structure. Simulation results for I_{out} vs V_{ds} curve for Wilson Current Mirror are shown in fig 6.

Advantages:

1. Curve is much flatter than basic and cascode current mirrors.
2. Output resistance becomes even much higher than cascode current mirror. This is caused by two positive feedback effects.

Disadvantages:

1. Current becomes constant for quite large value of V_{ds} e.g. in this case minimum V_{ds} is 1.22V.

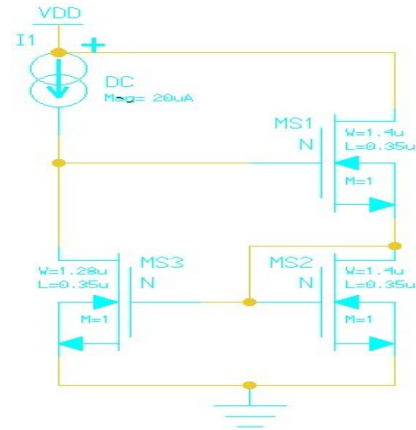


Fig 5 : Wilson current mirror

Simulation results:

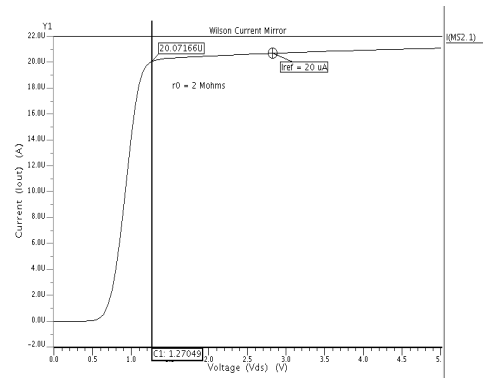


Fig 6 : I_{out} vs V_{ds} curve for Wilson Current Mirror.

V. LOW VOLTAGE SELF CASCODE CURRENT MIRROR

A self cascode current mirror is proposed that required a low bias voltage of order of $\pm 1.0V$ [9, 10]. The selection criterion for I_3 is to ensure lower V_{in} . I_2 is selected to ensure ON condition for M6 (Fig:7). The aspect ratios of different transistors are given in TABLE1. The small signal transfer analysis of this circuit at $20 \mu A$ gave the current gain, i.e. $I_{out}/I_{in} = 1$, and output resistance as $10 M\Omega$. The power dissipation for this is high. Simulation results for I_{out} vs V_{ds} curve for Self Cascode Current Mirror are shown in fig 8. This approach of increasing the (W/L) aspect ratios works

effectively at low bias voltage V_{in} of 1 V making it quite attractive for biasing analog circuits requiring high output resistance and gain. Hence they can be used as load resistances in CM circuits. They can extensively be used where power supply requirements are not the constraint.

Advantages:

1. High performance since output current is constant for low value of V_{ds} .
2. High output impedance.

Disadvantages:

1. Power dissipation is high.

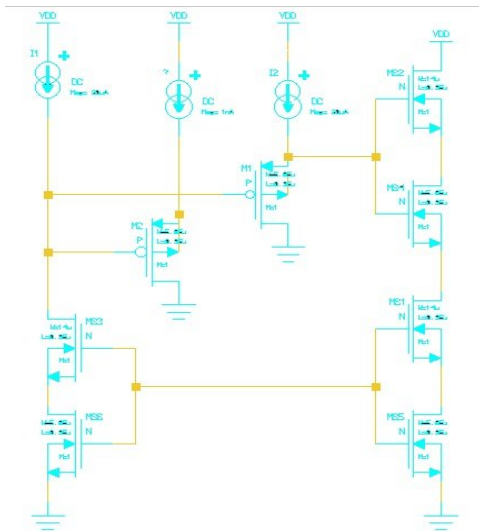


Fig 7 : Low voltage self cascode current mirror

$$\begin{aligned} I_1, I_2 &= 20 \text{ nA} \\ I_3 &= 1 \text{ nA} \\ V_2 &= 1 \text{ V} \end{aligned}$$

Design specifications:

MOSFETs	Type	W/L
MS1,MS2,MS3	NMOS	70 to 14/0.35
MS4,MS5,MS6	NMOS	5.25/0.35
M1,M2	PMOS	5.25/0.35

Table 1 : aspect ratios of all MOSFETS

Simulation results:

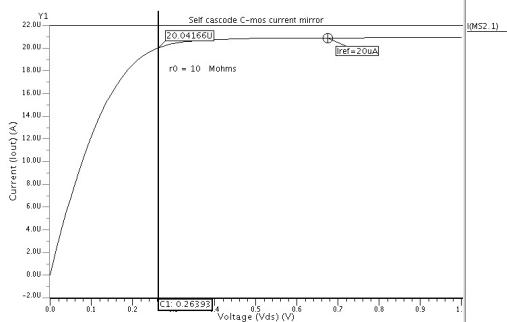


Fig 8 : I_{out} vs V_{ds} curve for Low voltage self cascode Current Mirror.

Comparison of different current mirrors:

A comparison of different current mirrors based on above simulation is given in TABLE 3. This TABLE compares the values of output impedance for each mirror and the minimum output voltage required for running the circuit.

Current Mirrors	Stability	Output resistance	Min. output voltage(V)
Basic current mirror	Poor	126 K	0.254
Cascode current mirror	Good	1.07 M	1.22
Wilson current mirror	Better	2 M	1.27
Low voltage	Excellent	10 M	0.26

Table 3 : Comparison of different current mirrors

REFERENCES REFERENCES REFERENCIAS

1. Behzad Razavi, "Design of Analog CMOS Integrated circuits", TMH edition 2002.
2. Yan Shouli, Sanchez-Sinencio Edgar, "Low voltage Analog Circuit Design Techniques", *IEICE Transactions: Analog Integrated Circuits and Systems*, vol EOO +A, no.2, pp1-3, February 2000.
3. G.Givstolisi, M. Consicione and F. Cutri, "A low-voltage low power voltage reference based on sub-threshold MOSFETs", *IEEE J. Solidstate circuits* 38, pp- 151-4, 2003.
4. S.S. Rajput and S.S. Jamuar, "Low voltage analog circuit design techniques", *IEEE circuits systems Magazine* 2, pp- 24- 42, 2002.
5. Alfonso Rincon M. Gabriel, "Low Voltage Design Techniques and Considerations for Integrated Operational Amplifier Circuit", *Georgia Institute of Technology, Atlanta, GA, 30332*, May 31, 1995.
6. T. Itakura and Z. Czarnul, "High output resistance CMOS current mirrors for low voltage applications." *IEICE Trans. Fundamentals*, vol. E80-A, no. 1, pp. 230-232, 1997.
7. J. Mulder, A.C. Woerd, W.A. Serdijn, and A.H.M. Roermund, "High swing cascode MOS current mirror." *Electron. Lett.*, vol. 32, pp. 1251- 1252, 1996.
8. E. Sackinger and W. Guggenbuhl, "A high swing, high impedance MOS cascode current mirror." *IEEE J. Solid State Circuits*, vol. 25, pp. 289- 298, 1990.
9. S.S. Rajput, "Low voltage current mode circuit structures and their applications." Ph.D. Thesis, Indian Institute of Technology, Delhi, 2002.
10. S.S. Rajput and S.S. Jamuar, "A current mirror for low voltage, high performance Analog Circuits." In *Proc. Analog integrated Circuits & Signal, Kluwer Academic Publications*, 36, pp. 221-233, 2003.
11. Xie, M.C. Schneider, E. Sanchez-Sinencio, and S.H.K. Embabi, "Sound design of low power nested transconductance – capacitance compensation amplifiers," *Electronic letters*, Vol. 35, pp 956-958, June 1999.



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Towards The Solution of Variants of Vehicle Routing Problem

By Pawan Jindal

Jaypee University of Engineering & Technology, Guna, M.P. India

Abstract - Some of the problems that are used extensively in real life are NP complete problems. There is no any algorithm which can give the optimal solution to NP complete problems in the polynomial time in the worst case. So researchers are applying their best efforts to design the approximation algorithms for these NP complete problems. Approximation algorithm gives the solution of a particular problem, which is close to the optimal solution of that problem. In this paper, a study on variants of vehicle routing problem is being done along with the difference in the approximation ratios of different approximation algorithms as being given by researchers and it is found that Researchers are continuously applying their best efforts to design new approximation algorithms which have better approximation ratio as compared to the previously existing algorithms.

Keywords : *Approximation algorithms, Vehicle Routing problem with time windows, NP completeness.*

GJCST Classification : C.2.2



Strictly as per the compliance and regulations of:



Towards The Solution of Variants of Vehicle Routing Problem

Pawan Jindal

Abstract - Some of the problems that are used extensively in real life are NP complete problems. There is no any algorithm which can give the optimal solution to NP complete problems in the polynomial time in the worst case. So researchers are applying their best efforts to design the approximation algorithms for these NP complete problems. Approximation algorithm gives the solution of a particular problem, which is close to the optimal solution of that problem. In this paper, a study on variants of vehicle routing problem is being done along with the difference in the approximation ratios of different approximation algorithms as being given by researchers and it is found that Researchers are continuously applying their best efforts to design new approximation algorithms which have better approximation ratio as compared to the previously existing algorithms.

Keywords : Approximation algorithms, Vehicle Routing problem with time windows, NP completeness.

1. INTRODUCTION

Transportation supports most of the social and economic activities. The annual cost of excess travel in U.S.A. has been estimated approximately 45 billion USD [68] and the turnover of transportation of goods in Europe is approximately 168 billion USD per year. In United Kingdom, France and Denmark transportation represents approximately 15%, 9% and 15% of national expenditures respectively [70]. It is well known fact that the vehicle routing problem is a combinatorial optimization and integer programming problem in which service to finite number of customers with a fleet of vehicles is being done. Vehicle routing problem was proposed by Dantzig and Ramser in 1959. Vehicle routing problem is an important optimization problem in the fields of distribution, transportation, and logistics. In Vehicle Routing problem, goods have to be delivered to the customers who have placed orders for such goods in such a way so that the total cost of the delivering of goods to the customers can be minimized. In VRP each and every customer has a given demand and no any vehicles can service more customers than its predefined capacity. Many algorithms have been designed by researchers for searching for good solutions to the problem, but no any polynomial time algorithms have been designed which can give the exact solution of Vehicle Routing problem with time windows in the polynomial time in the worst case. There are several variations of the vehicle routing problems. In

Vehicle Routing Problem with Pickup and Delivery, a large number of goods have to be moved from certain pickup locations to other delivery locations and the goal is to find optimal routes for vehicles to visit the pickup and drop-off locations so that the cost of delivering the goods to the different locations can be minimized. In Vehicle Routing Problem with LIFO principle, a large number of goods have to be moved from certain pickup locations to other delivery locations and the goal is to find optimal routes for vehicles to visit the pickup and drop-off locations so that the cost of delivering the goods to the different locations can be minimized with the restriction of the item being delivered must be that item most recently picked up. The benefit of this scheme over the previous scheme is the reduction of the loading and unloading times at delivery locations because there is no need to temporarily unload items. In Vehicle Routing Problem with Time Windows, there are n numbers of cities. Each and every city has a particular time windows $[R(v), D(v)]$. Where $R(v)$ represents the releasing time for a particular vertex and $D(v)$ represents the deadline for a particular vertex and the goal is to visit the maximum number of cities with in their time windows. Time windows are when they can be considered non bidding for penalty cost. Time windows are called hard when they cannot be violated, i.e. if any vehicle reaches to a particular city too early so it must wait unless and until the time windows opens and the vehicle is not allowed to arrive late. In Capacitated Vehicle Routing Problem, the vehicles have limited capacity of the goods that must be delivered. In Capacitated Vehicle Routing Problem with time windows $[R(v), D(v)]$, the vehicles have limited capacity of the goods that must be delivered. The most important application of VPRTW includes deliveries to supermarkets, industrial refuse collection, routing of school bus, security patrol services, urban newspaper distribution etc. The Vehicle Routing Problem with Time Windows has already been studied in the literature of Operations Research ([1, 10]). Different heuristics [9, 17, 18, 19] like Simulated Annealing, local search, Genetic algorithms, cutting plane and branch and bound methods [20, 14, 16] have been proposed to get the optimal solution of this problem. For general graphs, when there are a constant number of different time windows Chekuri and Kumar [8] gave a constant-factor approximation to solve vehicle routing problem with time windows. Capacitated Vehicle Routing Problem (CVRP) is the homogeneous VRP. In Multiple Depot VRP

Author : Department of Computer Science & Engineering, Jaypee University of Engineering & Technology, Guna, M.P. India.
E-mail : pawan_intelligent88@yahoo.co.in

(MDVRP) customers get their deliveries from several depots. In VRP with Time Windows [71] each and every customer has time window $(R[v], D[v])$ where $D[v]$ represents the deadline for a particular customer while $R[v]$ represents the releasing time for a particular customer and the goal is to visit the maximum number of customers in their time windows $(R[v], D[v])$. In Stochastic VRP (SVRP) any customer may have a random behavior. In Periodic VRP (PVRP) delivery to the customer is being done in some days. In Split Delivery VRP (SDVRP) several vehicles serve a customer. In the split delivery vehicle routing problem (SDVRP) there is no any restriction of visiting each and every customer exactly visited once.

Additionally, the demand of each and every customer may be greater than the capacity of the vehicles. It is well known fact that the SDVRP is NP-hard problem, even under restricted conditions on the costs, when each and every vehicle have a capacity greater than two. But it can be solved in the time complexity of polynomial time when the vehicles have a maximum capacity of two. The cost saving that which can be obtained by allowing split deliveries can be up to 50% of the cost of the optimal solution of the VRP.

The variant of the VRP[71] in which the demand of a customer may be greater than the vehicle capacity, but vehicle has to serve every customer minimum number of the possible time. The cost saving which can be obtained by allowing more than the minimum number of required visits to each and every customer to be served by vehicle can be again up to 50%. Simple heuristics that serve the customers with demands greater than the vehicle capacity by full load out-and-back trips until the demands become less than the vehicle capacity may be quite far from the optimal solution.

Three heuristic methods [71] have been already proposed for the solution of the SDVRP: The local search, a simple and effective tabu search algorithm and a sophisticated heuristic which uses the information collected during the tabu search, builds promising routes and solves MILP models to decide which routes to use and how to serve the customers through those routes to obtain the solution which is close to the optimal solution. The heuristics will be compared on a set of benchmark instances.

In VRP with Backhauls (VRPB) vehicle must pick something up from the customer after all deliveries are done to the customers. VRP with Backhauls (VRPB) is also known as the linehaul-backhaul problem which is an extension of the Capacitated VRP (CVRP) where the customer set is partitioned into two subsets. The first subset contains the linehaul customers; each requires a given quantity of product which has to be delivered. The second subset contains the backhaul customers, where a given quantity of inbound product must be picked up. This customer partition is extremely frequent in practical

situations. Grocery industry is a common example, where supermarkets and shops are the linehaul customers and grocery suppliers are the backhaul customers. It has been widely recognized that in this mixed distribution—collection context a significant saving in transportation costs can be achieved by visiting backhaul customers in distribution routes. More precisely, the VRPB [71] can be stated as the problem of determining a set of vehicle routes visiting all customers, and (a) each vehicle performs exactly one route; (b) each route starts as well as finishes at the depot; (c) for each route the total load associated with linehaul and backhaul customers should never exceed, separately, the vehicle capacity; (d) on each route the backhaul customers, are visited after all linehaul customers; and (e) the total distance traveled by the vehicles to serve the customers is minimized. The constraint (d) is practically motivated by the fact that vehicles are rear loaded which proves that the onboard load rearrangement required by a mixed service is difficult to carry out at customer locations. The most important reason is that, in many applications, line haul customers have a higher service priority as compared to backhaul customers. In VRP with Pick-Ups and Deliveries (VRPPD) the vehicle picks something up and delivers it to the customer.

II. VEHICLE ROUTING PROBLEM WITH TIME WINDOWS

If $V[72]$ represents set of vehicles and all vehicles are considered to be identical. C represents set of customers. $G=(N, A)$ represents directed graph where N represents the set of vertices of graph while E represents the set of edges of graph. This particular directed graph consists of $|C|+2$ number of vertices, where customers are denoted $1,2,3,\dots,n$ and the depot is represented by the vertex "0"(the starting depot) and the vertex " $n+1$ "(the returning depot). N is the set of vertices. There is no edge ending at the vertex "0" or originating from the vertex " $n+1$ ". c_{ij} (where $i \neq j$) represents the cost of traveling from the vertex i to the vertex j . t_{ij} (where $i \neq j$) represents the service time at the customer i . Each vehicle has a capacity q and each customer i has a demand d_i . Each and every customer has a time window $[a_i, b_i]$ and a vehicle must arrive at the customer before b_i . If any vehicle arrives to the customer before the time windows opens, that vehicle has to be wait until a_i to service the customer. The time windows for both depots are assumed to be identical to $[a_0, b_0]$ which represents the scheduling horizon. The vehicle cannot leave the depot before a_0 and must return at the latest at time b_{n+1} . It is assumed that $q, a_i, b_i, d_i, c_{ij}, t_{ij}$ are positive integers. It is also assumed that the triangle inequality [72] is satisfied for both c_{ij} and t_{ij} . There are two sets of decision variables x and s . For each edge (i,j) where $i \neq j, i \neq n+1, j \neq 0$ and each vehicle k we define x_{ijk} as

$$x_{ijk} = \begin{cases} 1, & \text{if vehicle } k \text{ drives directly from vertex } i \text{ to the vertex} \\ 0, & \text{otherwise.} \end{cases}$$

The decision variable s_{ik} is defined for each and vertex i and each vehicle k and denoted the time when the vehicle k starts to service the customer i . When the vehicle k does not service to the customer i , s_{ik} has no meaning and consequently its value is considered irrelevant. As we have assumed $a_0=0$ and therefore $s_{0k}=0$ for all k . The goal in the case of VRPTW is to design a set of routes that minimizes the total cost, such that 1, if vehicle k drives directly from vertex i to the vertex

- (a) Each customer is being served by vehicle exactly once.
- (b) Every route starts at the vertex 0 and ends at vertex $n+1$.
- (c) The time windows of the customers and capacity constraints of the vehicles are being observed carefully.

The above informal definition of VRPTW can be stated mathematically as a multicommodity network flow problem with time windows and the capacity constraints:

$$\min \sum_{k \in V} \sum_{i \in N} \sum_{j \in N} c_{ij} x_{ijk} \quad \text{such that} \quad (1)$$

$$\sum_{k \in V} \sum_{j \in N} x_{ijk} = 1 \quad \forall i \in C \quad (2)$$

$$\sum_{i \in C} d_i \sum_{j \in N} x_{ijk} \leq q \quad \forall k \in V \quad (3)$$

$$\sum_{j \in N} x_{0jk} = 1 \quad \forall k \in V \quad (4)$$

$$\sum_{i \in N} x_{ihk} - \sum_{j \in N} x_{jhk} = 0 \quad \forall h \in C, \forall k \in V, \quad (5)$$

$$\sum_{j \in N} x_{i,n+1,k} = 1 \quad \forall k \in V \quad (6)$$

$$x_{ijk} (s_{ik} + t_{ij} - s_{jk}) \leq 0 \quad \forall i, j \in N, \forall k \in V \quad (7)$$

$$a_i \leq s_{ik} \leq b_i \quad \forall i \in N, \forall k \in V \quad (8)$$

$$x_{ijk} \in \{0, 1\} \quad \forall i, j \in N, \forall k \in V \quad (9)$$

The main goal of the objective function (1) is to minimize the total travel cost. The constraint (2) ensures that each customer is visited exactly once while constraint (3) ensures that a vehicle can only be loaded up to its capacity. Equations 4 indicated that each and every vehicle must leave the depot 0. Equation 5

indicates that when a vehicle arrives at a customer it must leave for another destination. Equation 6 indicates that all vehicles must arrive at the depot $n+1$. Inequality (7) indicates the relationship between the departure time of vehicle from the customer and its immediate successor. Constraint (8) indicates the observation of time windows. Integrality constraints are shown by (9). The model for the representation of VRPTW can also incorporate a constraint giving an upper bound on the number of vehicles, as is the case in Desrosiers, Dumas, Solomon and Soumis [59].

III. VEHICLE ROUTING PROBLEM WITH TIME WINDOWS IS NP COMPLETE PROBLEM

Sorting algorithms like selection sort, bubble sort, insertion sort are known as quadratic sorting because these algorithms have time complexity of $O(n^2)$ in the worst case where n is the size of input. Sorting algorithms like counting sort, radix sort and bucket sort have linear time complexity in the worst case. So these algorithms are known as sorting in linear time. It is well known fact that all problems cannot be solved in polynomial time in the worst case. For example, Turing's famous "Halting Problem," which cannot be solved by any computer, no matter how much time is provided. Those problems that can be solved, but in time $O(n^k)$ for any constant k are known as tractable problems or easy problems. For example sorting algorithms like selection sort, bubble sort, insertion sort, counting sort, radix sort and bucket sort can give sorted output in the time complexity of polynomial time so sorting problems are tractable problems or easy problems. Those problems that can not be solved, in polynomial time $O(n^k)$ for any constant k but they require super-polynomial time for their executions are known as intractable problems or hard problems. No polynomial-time algorithm has yet been discovered for an NP-complete problem which can solve the problem in the polynomial time in the worst case nor anyone has been able to prove that no polynomial-time algorithm can exist for any one of NP-complete problems. So $P \neq NP$ question has been one of the deepest, most perplexing open research problems in theoretical computer science since it was first proposed in 1971.

The class P consists of those problems which can be solved in polynomial time in the worst case. More specifically, they are problems that can be solved in time $O(n^k)$ for some constant k , where n is the size of the input to the problem. Problems like sorting problem, searching problems are in class P . The class NP consists of those problems that are "verifiable" in polynomial time. If a "certificate" of a solution is being given, then we could verify that the certificate is correct in time polynomial in the size of the input to the problem. Hamiltonian-cycle problem is in class NP because In Hamiltonian-cycle problem, directed graph $G = (V, E)$ is being given, and then certificate of this problem would

be a sequence $v_1, v_2, v_3, \dots, v_n$ of $|V|$ vertices. It is easy to check in polynomial time whether these set of vertices would be lead to the Hamiltonian cycle or not. Also, 3-CNF satisfiability problem is in class NP. It is well known fact that any problem in class P will also in class NP, because if a problem is in P then we can solve it in polynomial time without even being given a certificate. Any problem will be lie in the class NPC-and we refer to it as being NP-complete-if the problem is in NP and is as "hard" as any problem in NP. The first NP complete problem is the circuit-satisfiability problem, in which we are given a Boolean combinational circuit which is being consists of AND, OR, and NOT gates and the question is to know whether there is any set of Boolean inputs to this circuit that causes its output to be 1. It is well known fact that the concept of NP-complete was firstly introduced by Stephen Cook in 1971 in a paper entitled "The complexity of theorem-proving procedures" on pages 151-158 of the Proceedings of the 3rd Annual ACM Symposium on Theory of Computing, Although the term NP-complete did not appear anywhere in his paper. At that conference, there was a debate among the computer researchers about whether NP-complete problems could be solved in polynomial time on a deterministic Turing machine or not. At that time, John Hopcroft brought everyone at the conference to a consensus that the question of whether NP-complete problems are solvable in polynomial time or not must be put off to be solved at some later time, since nobody had any formal proofs for their claims. No any scientist has yet been able to prove conclusively whether NP-complete problems are solvable in polynomial time or not. Also The Clay Mathematics Institute is offering a US\$1 million reward to any researcher who has a formal proof that $P=NP$ or that $P \neq NP$. Researchers are continuously doing hard work in this field to give the formal prove of either $P=NP$ or $P \neq NP$ but did not achieve success in this field till date. Also In the celebrated Cook-Levin theorem, Cook proved that the Boolean satisfiability problem is NP-complete problem. In 1972, Richard Karp proved that several other problems are also NP-complete; So it shows that there is a class of NP-complete problems. Satisfiability, 0-1 Integer Programming, Clique, Set Cover, Vertex Cover, Set Covering, Feedback Node Set, Feedback Arc Set, Directed Hamilton Circuit Undirected Hamilton Circuit, Satisfiability With At Most 3 Literals Per Clause, Chromatic Number, Cover, Exact, Hitting, Steiner Tree, 3-Dimensional Matching Knapsack (Karp's definition of Knapsack is closer to Subset sum), Job Sequencing, Partition Max Cut. After then, thousands of other problems have been shown to be NP-complete problems by reductions from other problems previously shown to be NP-complete. There is no any algorithm which can solve Vehicle Routing Problem with time windows in the polynomial time in the worst case. It is well known fact that Vehicle Routing Problem with time

windows is NP complete problem. So researchers are continuously applying their best efforts to give approximation algorithm for Vehicle Routing Problem with time windows. As it well known fact that approximation algorithms give the approximation result which is close to the optimal value to a particular problem. Most of the problems of practical significance are NP-complete but we cannot avoid them. There are three approaches to getting around NP-completeness of the problems. First, if the size of inputs is small, an algorithm which has exponential running time may be a satisfactory algorithm to solve a problem. Second, we may isolate those special cases that are solvable in polynomial time. Third, we can find out the solution which is close to the optimal solution of the problem and that solution can be easily found out with the help of the polynomial time approximation algorithm. Suppose there is an optimization problem in which each potential solution has a positive cost. It is well known fact that the optimization problem can be divided in the two parts. The maximization problem and the minimization problem. In maximization problem, we want to maximize the value of output to a problem and in minimization problem we want to minimize the value of output. If the size of the input of a problem is n . Let C^* be the cost as being obtained by optimal solution of a problem and C is the cost as being obtained by approximation algorithm [49]. Then the approximation ratio $\rho(n)$ is defined as

$$\text{Maximum } (C/C^*, C^*/C) \leq \rho(n).$$

The definitions of approximation ratio and of $\rho(n)$ -approximation algorithm apply for both minimization and maximization problems. It is well known fact that for a maximization problem, $0 < C \leq C^*$ and the ratio C^*/C gives the ratio by which the cost of optimization algorithm is larger than that of the cost of the approximate solution. Also, for a minimization problem, $0 < C^* \leq C$, and the ratio C/C^* gives the ratio by which the cost of the approximate solution is greater than the cost of an optimal solution. Since all solutions are assumed to have positive cost, these ratios are always well defined. Also the approximation ratio of an approximation algorithm [49] is never less than 1. For some problems, there are polynomial-time approximation algorithms which have small constant approximation ratios, but for other problems, the best known polynomial-time approximation algorithms have approximation ratios that grow as functions of the size of input of the problem.

An approximation scheme for an optimization problem [49] is defined as an approximation algorithm which takes as input an instance of the problem, along with a value $\epsilon > 0$ such that for any fixed, the scheme is a $(1 + \epsilon)$ -approximation algorithm. An approximation Scheme is said to be a polynomial-time approximation scheme if and only if for any fixed $\epsilon > 0$, the scheme runs in polynomial time in the size n of its input instance.

IV. COMPARISON OF APPROXIMATION ALGORITHMS FOR VEHICLE ROUTING PROBLEM

Arkin, Mitchell and Narasimhan [26] gave first non-trivial $(2 + \varepsilon)$ approximation algorithm for orienteering points in the Euclidean plane. A. Blum, S. Chawla, D. Karger, T. Lane, A. Meyerson, and M. Minkoff [6] gave the first approximation algorithm with a ratio of 4 for points in arbitrary metric spaces. After then N. Bansal, A. Blum., S. Chawla, and A. Meyerson [24] designed a new approximation algorithm for orienteering problem which improved this ratio to 3. A related problem to the orienteering problem is the minimum excess problem as being defined in [6]. In [6] the pseudo code for the orienteering problem depends upon the pseudo code of the min-excess problem. Also the min-excess problem can be approximated using algorithms for the k-stroll problem. In the k-stroll problem, the goal is to find a minimum length walk from source vertex s to target vertex t that visits at least k vertices. It is well known fact that k-stroll problem and the orienteering problem are equivalent to each other in terms of exact solvability because in both of these problems, the mission is to find the minimum length path from the source vertex s to the destination vertex t which covers maximum number of distinct vertices. The results in [6, 24] are based on existing approximation algorithms for k-stroll in undirected graphs. N. Bansal, A. Blum, S. Chawla and A. Meyerson.[24] gave approximation algorithm for Deadline-TSP which has the time complexity of $O(\log n)$ and they gave approximation algorithm for Vehicle Routing problem with time windows which has time complexity of $O(\log^2 n)$. Further they gave a bicriteria approximation algorithm for Deadline-TSP as well as for Vehicle Routing problem with time windows. If $\varepsilon > 0$, their bicriteria approximation algorithm produces a $\log(1/\varepsilon)$ approximation, while deadlines exceeds by a factor of $(1 + \varepsilon)$. C. Chekuri, N. Korula, and M. Pal [42] designed $(2 + \varepsilon)$ approximation for orienteering in undirected graphs, which improves upon the 3-approximation of [24]. C. Chekuri, N. Korula, and M. Pal [42] designed an improved $O(\log^2 \text{OPT})$ approximation for orienteering in directed graphs, where $\text{OPT} \leq n$ is the number of vertices visited by an optimal solution which improves over the previously result. Further it was being proved that for the time-window problem, an $O(\log \text{OPT})$ approximation can be easily achieved even for directed graphs if the algorithm is allowed quasi-polynomial time. As it has been already discussed that the best known polynomial time approximation ratios for Vehicle Routing problem with time windows are $O(\log^2 \text{OPT})$ for undirected graphs and $O(\log^4 \text{OPT})$ in directed graphs. If $D(v)$ represents the deadline for a particular vertex v and $R(v)$ represents the releasing time for a particular vertex v . Let $L(v) = D(v) - R(v)$ denotes the length of the

time-window for the vertex v and let $L_{\max} = \max_v L(v)$ and $L_{\min} = \min_v L(v)$. Let a be the known approximation ratio for orienteering problem. As $a = O(\log^2 \text{OPT})$ for directed graphs. C. Chekuri and N. Korula. Designed an $O(\log L_{\max})$ approximation when $R(v)$ and $D(v)$ both are integer valued for each v and they designed an $O(a \max\{\log \text{OPT}, \log L_{\max}/L_{\min}\})$ approximation. They also designed an $O(\log L_{\max}/L_{\min})$ approximation when there is no starting vertex and terminating vertex is being defined. Early surveys of solution techniques for the VRPTW[67] can be found in Golden and Assad [57], Desrochers et al. [58], and Chiang & Russell[66]. Desrosiers et al. [59] and Cordeau et al.[60] gave exact solution techniques for VRPTW. The complete explanation of these exact techniques can be found in Larsen [61] and Cook and Rich [62]. Researchers designed different approximation algorithms for VRPTW based on different designing techniques like Dynamic programming, Simulated Annealing etc. Fleischmann [63] and Taillard et al.[64] have used heuristic for VRP without time windows. In Taillard et al. [64], have designed solutions to the classical vehicle routing problem by using a TS heuristic. The routes which are obtained combine to produce workdays for the vehicles by solving a bin packing problem, an idea which is previously introduced in Fleischmann [63]. Compbell and Savelsbergh [65] has reported about insertion heuristics which can efficiently handle different types of constraints including time windows and multiple uses of vehicles. Compbell and Savelsbergh [65] introduced the home delivery problem which is the variant of Vehicle Routing problem and it is more closely related to real-world applications. Current VRPTW heuristics can be categorized as follows: (i) construction heuristics, (ii) improvement heuristics and (iii) meta-heuristics. Construction heuristics are sequential or parallel algorithms which aims at designing initial solutions to routing problems that can be easily improved upon by meta-heuristics or improvement heuristics. Sequential algorithms are being used to build a route for each vehicle, one after another with the help of decision functions for the selection of the customer which has to be inserted in the route and the insertion position within the route. Parallel algorithms build the routes for all vehicles in parallel by using a pre-computed estimate of the number of routes.

V. CONCLUSIONS

In this paper, a study on variants of vehicle routing problem is being done along with the difference in the approximation ratios of approximation algorithm as being given by researchers and it is found that Researchers are continuously applying their best efforts to design new approximation algorithms which have better approximation ratio as compared to the previously existing approximation algorithms. Researchers are proposing new heuristics for variants of Vehicle Routing problems.

REFERENCES REFERENCES REFERENCIAS

1. E.M. Reingold, J. Nievergelt and N. Deo, "Combinatorial Algorithms: Theory and Practice", Prentice-Hall, 1977.
2. Allahverdi, J. Gupta and T. Aldowaisan. A review of scheduling research involving setup considerations. *Omega, Int. Journal of Management Science*, 27:219–239, 1999.
3. E. M. Arkin, J. S. B. Mitchell, and G. Narasimhan. Resource-constrained geometric network optimization. In *Symposium on Computational Geometry*, pages 307–316, 1998.
4. B. Awerbuch, Y. Azar, A. Blum, and S. Vempala. Improved approximation guarantees for minimum-weight k-trees and prizecollecting salesmen. *Siam J. Computing*, 28(1):254–262, 1999.
5. R. Bar-Yehuda, G. Even, and S. Shih. On approximating a geometric prize-collecting traveling salesman. In *Proceedings of the European Symposium on Algorithms*, 2003.
6. Blum, S. Chawla, D. Karger, T. Lane, A. Meyerson, and M. Minkoff. Approximation algorithms for orienteering and discounted-reward tsp. In *Proceedings of the 44th Foundations of Computer Science*, 2003.
7. J. Bruno and P. Downey. Complexity of task sequencing with deadlines, set-up times and changeover costs. *SIAM Journal on Computing*, 7(4):393–404, 1978.
8. G.N. Frederickson and B. Wittman. Approximation algorithms for the traveling repairman and speeding deliveryman problems with unit-time windows. *Proceedings of APPROX-RANDOM*, LNCS 4627:119–133, 2007.
9. M. Desrochers, J. Desrosiers, and M. Solomon. A new optimization algorithm for the vehicle routing problem with time windows. *Operations Research*, 40:342–354, 1992.
10. M. Desrochers, J. Lenstra, M. Savelsbergh, and F. Soumis. Vehicle routing with time windows: optimization and approximation. In B.L. Golden, A.A. Assad (eds.). *Vehicle Routing: Methods and Studies*, North-Holland, Amsterdam, pages 65–84, 1988.
11. B. Golden, L. Levy, and R. Vohra. The orienteering problem. *Naval Research Logistics*, 34:307–318, 1987.
12. M. Gravel, W. Price, and C. Gagn. Scheduling jobs in a alcan aluminium factory using a genetic algorithm. *International Journal of Production Research*, 38(13):3031–3041, 2000.
13. C. Jordan. A two-phase genetic algorithm to solve variants of the batch sequencing problem. *International Journal of Production Research (UK)*, 36(3):745–760, 1998.
14. M. Kantor and M. Rosenwein. The orienteering problem with time windows. *Journal of the Operational Research Society*, 43:629–635, 1992.
15. Y. Karuno and H. Nagamochi. A 2-approximation algorithm for the multi-vehicle scheduling problem on a path with release and handling times. In *Proceedings of the European Symposium on Algorithms*, pages 218–229, 2001.
16. Kolen, A. R. Kan, and H. Trienekens. Vehicle routing with time windows. *Operations Research*, 35:266–273, 1987.
17. S. Thangiah. Vehicle Routing with Time Windows using Genetic Algorithms, *Application Handbook of Genetic Algorithms: New Frontiers*, Volume II. Lance Chambers (Ed.). CRC Press, 1995.
18. K. Tan, L. Lee, and K. Zhu. Heuristic methods for vehicle routing problem with time windows. In *Proceedings of the 6th AI and Math*, 2000.
19. M. Savelsbergh. Local search for routing problems with time windows. *Annals of Operations Research*, 4:285–305, 1985.
20. S. R. Thangiah, I. H. Osman, R. Vinayagamoorthy and T. Sun. Algorithms for the vehicle routing problems with time deadlines. *American Journal of Mathematical and Management Sciences*, 13(3&4):323–355, 1993.
21. J. Tsitsiklis. Special cases of traveling salesman and repairman problems with time windows. *Networks*, 22:263–282, 1992.
22. J. Wisner and S. Siferd. A survey of u.s. manufacturing practices in make-to-order machine shops. *Production and Inventory Management Journal*, 1:1–7, 1995.
23. W. Yang and C. Liao. Survey of scheduling research involving setup times. *International Journal of Systems Science*, 30(2):143–155, 1999.
24. N. Bansal, A. Blum., S. Chawla, and A. Meyerson. Approximation algorithms for deadline-TSP and vehicle routing with time-windows. In *Proceedings of the 36th Annual ACM Symposium on Theory of Computing*, pages 166–174. ACM New York, NY, USA, 2004.
25. R.K. Ahuja, T.L. Magnanti, and J.B. Orlin. *Network Flows: Theory, Algorithms, and Applications*. Prentice Hall, Upper Saddle River, New Jersey, 1993.
26. E.M. Arkin, J.S.B. Mitchell, and G. Narasimhan. Resource-constrained geometric network optimization. In *Symposium on Computational Geometry*, pages 307–316, 1998.
27. S. Arora and G. Karakostas. A $2 + \epsilon$ approximation algorithm for the k- MST problem. *Mathematical Programming A*, 107(3):491–504, 2006. Preliminary version in *Proc. of ACM-SIAM SODA*, 754–759, 2000.
28. B. Awerbuch, Y. Azar, A. Blum, and S. Vempala. Improved approximation guarantees for minimumweight k-trees and prizecollecting

- salesmen. SIAM J. Computing, 28(1):254–262, 1998. Preliminary Version in Proc. of ACM STOC, 277–283, 1995.
29. E. Balas. The prize collecting traveling salesman problem. *Networks*, 19(6):621–636, 1989.
 30. A.M. Frieze, G. Galbiati, and F. Maffioli. On the worst-case performance of some algorithms for the asymmetric traveling salesman problem. *Networks*, 12(1):23–39, 1982.
 31. R. Bar-Yehuda, G. Even, and S. Shahar. On approximating a geometric prize-collecting traveling salesman problem with time windows. *Journal of Algorithms*, 55(1):76–92, 2005. Preliminary version in Proc. Of ESA, 55–66, 2003.
 32. N. Garg. A 3-approximation for the minimum tree spanning k vertices. In *Proceedings of the 37th Annual Symposium on Foundations of Computer Science*, pages 302–309. IEEE Computer Society, 1996.
 33. Blum, R. Ravi, and S. Vempala. A Constant-Factor Approximation Algorithm for the k -MST Problem. *Journal of Computer and System Sciences*, 58(1):101–108, 1999. Preliminary version in Proc. Of ACMSTOC, 1996.
 34. K. Chaudhuri, B. Godfrey, S. Rao, and K. Talwar. Paths, trees, and minimum latency tours. In *44th Annual Symposium on Foundations of Computer Science*, pages 36–45. IEEE Computer Society, 2003.
 35. C. Chekuri and A. Kumar. Maximum coverage problem with group budget constraints and applications. *Proceedings of APPROXRANDOM*, pages 72–83, 2004.
 36. C. Chekuri and M. P’al. A recursive greedy algorithm for walks in directed graphs. In *Proceedings of the 46th Annual Symposium on Foundations of Computer Science*, pages 245–253. IEEE Computer Society, 2005.
 37. C. Chekuri and M. P’al. An $O(\log n)$ Approximation for the Asymmetric Traveling Salesman Path Problem *Theory of Computing*, 3:197–209, 2007. Preliminary version in Proc. of APPROX, 95–103, 2005.
 38. N. Garg. Saving an epsilon: a 2-approximation for the k -MST problem in graphs. In *Proceedings of the 37th Annual ACM Symposium on Theory of computing*, pages 396–402. ACM, 2005.
 39. M. X. Goemans and D. P. Williamson. A general approximation technique for constrained forest problems. *SIAM Journal on Computing*, 24:296–317, 1995.
 40. B.L. Golden, L. Levy, and R. Vohra. The orienteering problem. *Naval Research Logistics*, 34(3):307–318, 1987.
 41. G. Gutin and A.P. Punnen, editors. *The traveling salesman problem and its variations*. Springer, Berlin, 2002.
 42. C. Chekuri, N. Korula, and M. P’al. Improved Algorithms for Orienteering and Related Problems. In *Proc. of ACM-SIAM SODA*, January 2008.
 43. E.L. Lawler, A.H.G. Rinnooy Kan, J.K. Lenstra, and D.B. Shmoys, editors. *The Traveling salesman problem: a guided tour of combinatorial optimization*. John Wiley & Sons Inc, 1985.
 44. Nagarajan and R. Ravi. Poly-logarithmic approximation algorithms for directed vehicle routing problems. *Proceedings of APPROXRANDOM*, LNCS 4627:257–270, 2007.
 45. P. Toth and D. Vigo, editors. *The vehicle routing problem*. SIAM Monographs on Discrete Mathematics and Applications. Society for Industrial Mathematics, Philadelphia PA, 2001.
 46. J.N. Tsitsiklis. Special cases of traveling salesman and repairman problems with time windows. *Networks*, 22(3):263–282, 1992.
 47. K. Chen and S. Har-Peled. The orienteering problem in the plane revisited. *SIAM J. on Computing*, 38(1):385–397, 2008. Preliminary version in Proc. of ACM SoCG, 247–254, 2006.
 48. E. Horowitz and S. Sahni, "Fundamental of Computer Algorithms", Computer Science Press, 1982.
 49. T. Cormen, C. Leiserson and R. Rivest, "Introduction to Algorithms", MIT Press, 1991.
 50. A.V. Aho, J.E. Hopcroft and J.D. Ullman, "The Design and Analysis of Computer Algorithms", Addison-Wesley, 1974.
 51. A.V. Aho, J.E. Hopcroft and J.D. Ullman, "Data Structures and Algorithms", Addison-Wesley, 1984.
 52. S. Baase, "Computer Algorithms: Introduction to Design and Analysis", Addison-Wesley 1988.
 53. L. Banachowski, A. Kreczmar and W. Rytter, "Analysis of Algorithms and Data Structures", Addison-Wesley, 1991.
 54. R. Baeza-Yates, "Teaching Algorithms", IV Iberoamerican Congress on Computer Science Education, Canela, Brazil, July 1995.
 55. G. Tel, "Introduction to Distributed Algorithms" Cambridge Univ. Press 1995.
 56. G- Brassard, and P. Bratley, "Algorithmics: Theory and Practice", Prentice-Hall, 1988.
 57. Bruce L. Golden, Arjang A. Assad, Edward A. Wasil: Introduction. *Computers & OR* 13(2-3): 107-108 (1986).
 58. *Vehicle Routing with Time Windows: Optimization and Approximation* – Desrochers, Lenstra, et al. – 1988.
 59. *Time constrained routing and scheduling* – Desrosiers, Dumas, et al. – 1995.
 60. Cordeau, G. Laporte and A. Mercier. "A unified tabu search heuristic for vehicle routing problems with time windows." *Journal of the Operational Research Society*, 52:928–936 (2001).
 61. Larsen J. 1999. "Parallelization of the Vehicle Routing Problem with Time Windows" Ph.D. thesis, Institute of Mathematical Modelling, Technical University of Denmark, Lyngby, Denmark.

62. Cook W. & Rich J.L. 1999. "A Parallel Cutting-Plane Algorithm for the Vehicle Routing Problems with Time Windows." Working Paper, Department of Computational and Applied Mathematics, Rice University, Houston, U.S.A.
63. Fleischmann B. 1990. "The vehicle routing problem with multiple use of vehicles." Working Paper. Fachbereich Wirtschaftswissenschaften, Universität Hamburg, Germany.
64. Taillard E.D., Laporte G. & Gendreau M. 1996. "Vehicle routing with multiple use of vehicles" *Journal of the Operational Research Society* 47: 1065–1070.
65. Campbell A.M. & Savelsbergh, M. 2004. "Efficient insertion heuristics for vehicle routing and scheduling problems" *Transportation Science* 38(3), 369–378.
66. Chiang W. & Russell R.A. 1996. "Simulated annealing metaheuristics for the vehicle routing problem with time windows." *Annals of Operations Research* 63: 3-27.
67. Liong Choong Yeun, Wan Rosmanira Ismail, Khairuddin Omar2 & Mourad Zirour 2008. "Vehicle Routing Problem: Models and Solutions *Journal of Quality Measurement and Analysis*" Volume 4, Issue 1, pp. 205-218.
68. King, G.F. and C.F. Mast (1997), "Excess Travel: Causes, Extent and Consequences", *Transportation Research Record* 1111, 126–134.
69. Crainic, T. G. and G. Laporte (1997), "Planning Models for Freight Transportation", *European Journal of Operational Research* 97, 409–438.
70. Vehicle Routing Problem with Time Windows, Part I: Route Construction and Local Search Algorithms.
71. http://osiris.tuwien.ac.at/~wgarn/VehicleRouting/vehicl_routing.html
72. <http://alvarestech.com/temp/vrptw/Vehicle%20Routing%20Problem%20with%20Time%20Windows.pdf>



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
Volume 11 Issue 15 Version 1.0 September 2011
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Avoiding Loops and Packet Losses in ISP Networks

By A.Venkata Rohit , V.Sunil kumar, S.Mohankrishna, G.Siddhartha
GITAM University

Abstract - Even in well managed Large ISP networks, failures of links and routers are common. Due to these failures the routers update their routing tables. Transient loops can occur in the networks when the routers adapt their forwarding tables. In this paper, a new approach is proposed that lets the network converge to its optimal state without loops and the related packet lossless. The mechanism (OUTFC-Ordered Updating Technique with Fast Convergence) is based on an ordering of the updates of the forwarding tables of the routers and fast convergence. Typically we have chosen a Network consisting of routers and Link costs for simulation. Link failures are simulated. Avoiding transient loops in each case is demonstrated, by constructing a Reverse Shortest PathTree (RSPT).

Keywords : *ISP networks, OUTFC (Ordered Updating Technique with Fast Convergence), Link Failures, Reverse Shortest Path Tree (RSPT).*

GJCST Classification : *C.2.1, C.2.2*



Strictly as per the compliance and regulations of:



Avoiding Loops and Packet Losses in ISP Networks

A.Venkata Rohit^α, V.Sunil kumar^Ω, S.Mohankrishna^β, G.Siddhartha^ψ

Abstract - Even in well managed Large ISP networks, failures of links and routers are common. Due to these failures the routers update their routing tables. Transient loops can occur in the networks when the routers adapt their forwarding tables. In this paper, a new approach is proposed that lets the network converge to its optimal state without loops and the related packet lossless. The mechanism (OUTFC-Ordered Updating Technique with Fast Convergence) is based on an ordering of the updates of the forwarding tables of the routers and fast convergence. Typically we have chosen a Network consisting of routers and Link costs for simulation. Link failures are simulated. Avoiding transient loops in each case is demonstrated, by constructing a Reverse Shortest Path Tree (RSPT).

Keywords : *ISP networks, OUTFC (Ordered Updating Technique with Fast Convergence), Link Failures, Reverse Shortest Path Tree (RSPT).*

I. INTRODUCTION

The link-state intra domain routing protocols that are used in ISP network [1] [2], were designed when IP networks were research networks carrying best-effort packets. The same protocols are now used in large commercial LSPs with stringent Service Level Agreements (SLA). Furthermore, for most Internet Service Providers, fast convergence in case of failures is a key problem that must be solved. Today, customers are requiring 99.99% reliability or better and providers try to avoid all packet losses.

Transient loops can be occurred due the topological change in the network when links failure occurred in network [1]. A network typically contains point-to-point links and LAN Point-to-point links are typically used between Points of Presence (POPs) while LANs are mainly used inside POPs. When a point-to-point link fails, two cases are possible. If the link is not locally protected, the IGP should converge as quickly as possible. Another source of changes in IP networks are the IGP metrics. Today, network operators often change IGP metrics manually to reroute some traffic in case of sudden traffic increase. Second type of important events is those that affect routers. Routers can fail abruptly, but often routers need to be rebooted for software upgrades.

To avoid transient loops during the convergence of link-state protocols, we propose to force

the routers to update their FIB by respecting an ordering that will ensure the consistency of the FIB of the routers during the whole convergence phase of the network [1]. In the context of a predictable maintenance operation, the resources undergoing the maintenance will be kept up until the routers have updated their FIB and no longer use the links to forward packets. In the case of a sudden failure of a link that is protected with a Fast Reroute technique, the proposed ordering ensures that a packet entering the network will either follow a consistent path to its destination by avoiding the failed component or reach the router adjacent to the failure and will be deviated by the Fast Reroute technique to a node that is not affected by the failure, so that it will finally reach its destination.

II. OUR APPROACH

Studies on the occurrence of failures in a backbone network have shown that failures of links and routers are common even in a well managed network [1]. On the other hand, an increasing number of users and services are relying on the Internet and expecting it to be always available. In order to ensure high availability in spite of failures, a routing scheme needs to quickly restore forwarding to affected destinations. Traditional routing schemes such as OSPF trigger link state advertisements in response to a change in topology, and cause network-wide recomputation of routing tables. Such a global rerouting incurs some delay before traffic forwarding can resume on alternate paths. During this convergence delay, routers may have inconsistent views of the network, resulting in forwarding loops and dropped packets [2].

OUTCF [7] was recently proposed to address the above concerns and achieves three interconnected objectives: 1) loop- free forwarding; 2) minimal convergence delay. At no time can a forwarding loop happen with OUTCF in the case of a single failure. OUTCF also reduces the period of disruption when packets are dropped due to the lack of valid routes. Lastly, OUTCF minimizes the convergence delay, i.e., packets are forwarded along optimal paths and the network is ready to absorb another change as soon as possible. The drawback of OUTCF, however, is that it requires each packet to carry the cost of the remaining path to the destination, which needs multiple bytes in the header. Our objective is to minimize this overhead while maintaining the benefits of OUTCF.

*Author ^{αβ} : Department of IT, E-mails : avrohit2000@gmail.com, sunil.veee@gmail.com, smkrishna@ieee.org,
Author ^ψ : Department of CSE, GITAM University
E-mail : siddhartha.gundapaneni@gmail.com*

III. RELATED WORK

The problem of avoiding transient loops during IGP convergence has rarely been studied in the literature although many authors have proposed solutions to provide loop-free routing. An existing approach to loop-free rerouting in a link state IGP [8] requires that the rerouting routers take care of routing consistency for each of their compromised destinations, separately. In fact, those mechanisms were inspired by distance-vector protocols providing a transiently loop-free convergence [7]. With this kind of approach, a router should ask and wait clearance from its neighbours for each destination for which it has to reroute. This implies a potentially large number of messages exchanged between routers, when many destinations are impacted by the failure. Every time a router receives clearance from its neighbours for a given destination, it can only update forwarding information for this particular one. This solution would not fit well in a Tier-1 ISP topology where many destinations can be impacted by a single topological change. Indeed, in such networks, it is common to have a few thousands of prefixes advertised in the IGP [5]. Note that those solutions do not consider the problem of traffic loss in the case of a planned link shutdown. In [6], a new type of routing protocol allowing improving the resilience of IP networks was proposed. This solution imposes some restrictions on the network topology and expensive computations on the routers. Moreover, they do not address the transient issues that occur during the convergence of their routing protocol. In [4], extensions to link-state routing protocols are proposed to distribute link state packets to a subset of the routers after a failure. This fastens the IGP convergence, but does not solve the transient routing problems and may cause suboptimal routing.

In [2][3], transient loops are avoided when possible by using distinct FIB states in each interface of the routers. Upon a link failure, the network does not converge to the shortest paths. Based on the new topology. Indeed, the failure is not reported. Instead, the routers adjacent to the failed link forward packets along alternate links, and other routers are prepared to forward packets arriving from an unusual interface in a consistent fashion towards the destination. As such, the solution is a Fast Reroute technique. Our solution is orthogonal to [9] as our goal is to let the network actually converge to its optimal forwarding state by avoiding transient forwarding loops when a Fast Reroute mechanism has been activated, or when the failure is planned.

IV. METHOD TO HANDLE LOOPS

Each router will maintain one waiting list associated with each link being shut down during the RSPT computations. A rerouting router R will update its

FIB for a destination (which means that its paths to contain one or more links of the SRLG) once it has received the completion messages that unlock the FIB update in for one of the links being shut down. When updating its FIB, selects the outgoing interfaces for destination according to the new topology, i.e., by considering the removal or the metric increase of all the affected links. The meaning of a completion message concerning a link sent by a router is that has updated its FIB for all the destinations that it was reaching via before the event[8]. Let us now show that if a packet with destination reaches a rerouting router that has not performed its FIB update for destination, then all the routers on its paths to cannot have performed a FIB update for. If has not updated its FIB for destination, it cannot have sent a completion message for any of the failing links that it uses to reach. The failing links that a router on uses to reach are used by to reach, so that cannot have received all the necessary completion messages for any of those links. In other words, did not send a completion message for the links that it uses to reach. Thus, locks the FIB update for those links long its paths towards them. We provide the pseudo code that implements the ordering with completion messages. To process the metric increase (or shutdown) of a set of link , a router will compute the reverse SPT rooted on each link belonging to , that it uses in its current, outdated SPT. During this computation, it will obtain the rank associated with. It will then record the next-hops that it uses to reach in a list. These are the neighbors to which it will send a completion message concerning link. If the rank associated with a link is equal to zero, then updates its FIB directly for the destinations that it reaches via this link, and it sends a completion message to the corresponding next-hops. In the other cases, builds the waiting list associated with, containing the neighbors that are using to reach, and it starts the timer considering the rank associated with this link[9]. Once a waiting list for a link becomes empty or its associated timer elapses, can update its FIB for all the destinations that it reached via this link and send its own completion message towards the neighbors that it used to reach the link.

// Computation of the RSPTs of the affected link used by R

for each Link $X \rightarrow Y \in S$ do

if $X \rightarrow Y \in \text{SPTold}(R)$ then

//Computation of the rSPT

Link RSPT = rSPT ($X \rightarrow Y$);

//Computation of the rank

LinkRank = depth(R, Link RSPT);

//Computation of the set of neighbors to which a

//Completion message concerning this link will be sent

$I(X \rightarrow Y) = \text{NextHops}(R, X \rightarrow Y)$;

if LinkRank == 0 then

// R is a leaf in rSPT ($X \rightarrow Y$),

```

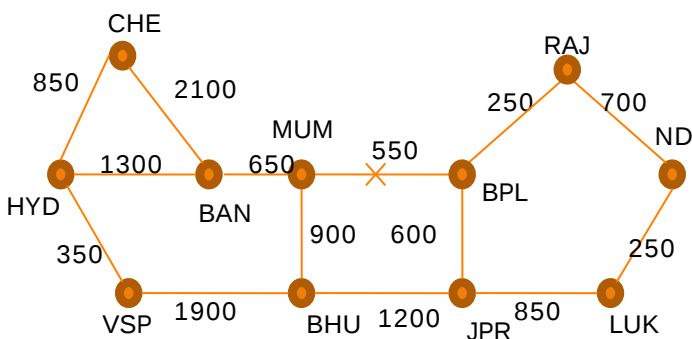
// it can updates its FIB directly
foreach d:  $X \rightarrow Y \in \text{Pathold}(R, d)$  do
    UpdateFIB(d);
end
//R can send its completion message for this link
foreach  $N \in I(X \rightarrow Y)$  do
    send(N, CM( $X \rightarrow Y$ ));
end
end
else
    //R is not a leaf in  $r\text{SPT}(X \rightarrow Y)$ ,
    //Computation of the waiting list
    WaitingList( $X \rightarrow Y$ ) = Childs(R, LinkRSPT);
    //Start the timer associated with this link.
    StartTimer( $X \rightarrow Y$ , LinkRank * MAXFIBTIME);
end
end
end

Upon reception of CM( $X \rightarrow Y$ ) from Neighbor N:
WaitingList( $X \rightarrow Y$ ).remove(N);
Upon (WaitingList( $X \rightarrow Y$ ).becomesEmpty())
Timer( $X \rightarrow Y$ ).hasExpreide();

//All the necessary completing message have been
received for
//The link or the timer associated with this link has
expired
//Update the FIB for each destination that was reached
via this link
foreach d:  $X \rightarrow Y \in \text{Path}(R, d)$  do
    UpdateFIB(d);
end
//Send completion message to the neighbor
that were used to reach this link
foreach  $N \in I(X \rightarrow Y)$  do
    send(N, CM( $X \rightarrow Y$ ));
end
end

Pseudo code for Avoiding Link Failures\
    
```

We consider a network to explain how to avoid the transient loops occur in the network by converging link state routing protocol. The Indian cities are connected in this net work like Mumbai (MUM), New Delhi(ND), Hyderabad(HYD), Madras etc.



Example : Internet topology with IGP

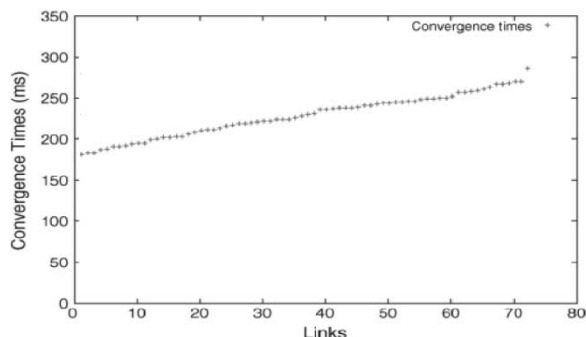
To understand this problem, let us consider the Internet2/Abilene backbone.¹ Fig. 1 shows the IGP topology of this network. Assume that the link between MUM and BPL fails but was protected by an MPLS tunnel between BPL and MUM via JPR and BHU. When JPR receives a packet with destination BAN, it forwards it to BPL, which forwards it back to JPR, but inside the protection tunnel, so that MUM will decapsulate the packet, and forwards it to its destination, BAN.

This suboptimal routing should not last long, and thus after a while the routers must converge, i.e., adapt to the new shortest paths inside the network, and remove the tunnel. As the link is protected, the reach ability of the destinations is still ensured and thus the adaptation to the topological change should be done by avoiding transient loops rather than by urging the updates on each router. The new LSP generated by BPL indicates that BPL is now only connected to RAJ and JPR. Before the failure, the shortest path from LUK to MUM, BAN, CHE and HYD was via ND, RAJ and BPL. After the failure, ND will send its packets to MUM, BAN, CHE and HYD via LUK, JPR and BHU. During the IGP convergence following the failure of link MUM–BPL, transient loops may occur between ND and LUK depending on the order of the forwarding table updates performed by the routers. If ND updates its FIB before LUK, the packets sent by ND to MUM via LUK will loop on the LUK–ND link. To avoid causing a transient loop between LUK and ND, LUK should update its FIB before ND for this particular failure. A detailed analysis of the Internet2 topology shows that transient routing loops may occur during the failure of most links, except CHE–BAN and CHE–HYD. The duration of each loop will depend on how and when the FIB of each router is updated. Measurements on commercial routers have shown that updating the FIB may require several hundred of milliseconds. Transient routing loops of hundred milliseconds or more are thus possible and have been measured in real networks. As shown with the simple example above, the transient routing loops depend on the ordering of the updates of the FIBs. In the remainder of this paper, This proof is constructive as we give an algorithm that routers can apply to compute the ranks that let them respect the proposed ordering.

V. CONVERGENCE TIMES IN ISP NETWORKS

In this section, we analyze by simulations the convergence time of the proposed technique, in the case of a link down event. The results obtained for link up events are very similar. Indeed, the updates that are performed in the FIB of each router for the shutdown of a link impact the same prefixes for the linkup of the link. The only difference in the case of a link up is that the routers do not need to compute a reverse Shortest Path Tree. As no packet are lost during the convergence process.

Lsp_process_delay	[2,4]ms
Update_hold_down	180ms
rspt_computation_tome	[3,5]ms
Completion_message_process_delay	[2,4]ms
Completion_message_sending_delay	[2,4]ms



We cannot define the convergence time as the time required bringing the network back to a consistent forwarding state, as it would always be equal to zero. What is interesting to evaluate here is the time required by the mechanism to update the FIB of all the routers by respecting the ordering.

VI. EXPERIMENTAL RESULTS

```
C:\ C:\Mcc\TC.EXE
Enter number of routers : 11

Enter link 1(0 0 to quit) : 1
Enter weight for this link : 850
Enter link 2(0 0 to quit) : 1
Enter weight for this link : 1300
Enter link 3(0 0 to quit) : 1
Enter weight for this link : 350
Enter link 4(0 0 to quit) : 2
Enter weight for this link : 2100
Enter link 5(0 0 to quit) : 3
Enter weight for this link : 650
Enter link 6(0 0 to quit) : 4
Enter weight for this link : 1900
```

```
C:\ C:\Mcc\TC.EXE
Enter link 13(0 0 to quit) : 9
Enter weight for this link : 700
Enter link 14(0 0 to quit) : 10
Enter weight for this link : 250
Enter link 15(0 0 to quit) : 0

The adjacency matrix is :
08501300350 0 0 0 0 0 0 0 0 0
850 02100 0 0 0 0 0 0 0 0 0
13002100 0 0650 0 0 0 0 0 0 0
350 0 0 0 01900 0 0 0 0 0 0
0 0650 0 0900550 0 0 0 0 0 0
0 0 0 01900900 0 01200 0 0 0 0
0 0 0 0550 0 0600250 0 0 0 0
0 0 0 0 0 01200600 0 0850 0
0 0 0 0 0 0250 0 0 0700 0
0 0 0 0 0 0 0850 0 0250 0
0 0 0 0 0 0 0 0700250 0 0

Enter source node(0 to quit) :
```

```
C:\ C:\Mcc\TC.EXE
0 0 0 0 0 0 0 0700250 0

Enter source node(0 to quit) : 10
Enter destination node(0 to quit) : 3
Shortest distance is : 2400
Shortest Path is : 10->11->9->7->5->3
Enter source node(0 to quit) : 8
Enter destination node(0 to quit) : 3
Shortest distance is : 1800
Shortest Path is : 8->7->5->3
Enter source node(0 to quit) : 11
Enter destination node(0 to quit) : 3
Shortest distance is : 2150
Shortest Path is : 11->9->7->5->3
Enter source node(0 to quit) : 9
Enter destination node(0 to quit) : 3
Shortest distance is : 1450
Shortest Path is : 9->7->5->3
Enter source node(0 to quit) : 7
Enter destination node(0 to quit) : 3
Shortest distance is : 1200
Shortest Path is : 7->5->3
Enter source node(0 to quit) : 0
Enter destination node(0 to quit) :
```

```
C:\ C:\Mcc\TC.EXE
Enter source node(0 to quit) : 0
Enter destination node(0 to quit) : 0

Enter failure link 1(0 0 to quit) : 5
Enter weight for this link : 0
Enter failure link 2(0 0 to quit) : 0

08501300350 0 0 0 0 0 0 0 0
850 02100 0 0 0 0 0 0 0 0
13002100 0 0650 0 0 0 0 0 0 0
350 0 0 0 01900 0 0 0 0 0 0
0 0650 0 0900 0 0 0 0 0 0
0 0 0 01900900 0 01200 0 0 0
0 0 0 0 0 0600250 0 0 0
0 0 0 0 0 01200600 0 0850 0
0 0 0 0 0 0250 0 0 0700 0
0 0 0 0 0 0 0850 0 0250 0
0 0 0 0 0 0 0 0700250 0 0

Enter source node(0 to quit) :
```

```
C:\ Turbo C++ IDE
Reverse Shortest Path tree are :
8->7
7->9
9->11
11->10
Weight of spanning tree is : 7050
Reverse Shortest Path tree are :
5->6
3->5
1->3
1->2
1->4

The proposed order is
2 4 6 8 10 1 3 5 7 11 9
```

VII. CONCLUSION

In the proposed work, we have initially described the various types of topological changes that can occur in large IP networks. When failures occurs in the network the routers updates routing tables. Those updates may cause transient loops and each loop may cause packet losses or delays. Large ISPs require

solutions to avoid transient loops after those non-urgent events. To protect the network from transient loops, we propose OUTFC method that it is useful to define an ordering on the updates of the FIBs. We have proposed an ordering applicable for the failures of protected links and the increase of a link metric and another ordering for the establishment of a new link or the decrease of a link metric. We have shown by simulations that our method avoids the loops and converges network to its optimal state.

REFERENCES REFERENCES REFERENCIAS

1. Avoiding Transient Loops during the Convergence of Link-State Routing Protocols Pierre Francois, Member, IEEE, and Olivier Bonaventure, Member, IEEE.
2. P. Francois and O. Bonaventure, "Avoiding transient loops during IGP Convergence in IP Networks," in *Proc. IEEE INFOCOM*, March 2005.
3. ISO, "Intermediate system to intermediate system routing information exchange protocol for use in conjunction with the protocol for providing the connectionless-mode network service (iso 8473)," ISO/IEC, Tech. Rep. 10589:2002, April 2002.
4. G. Iannaccone, C. Chuah, S. Bhattacharyya, and C. Diot, "Feasibility of IP restoration in a tier-1 backbone," *IEEE Network Magazine*, January-February 2004.
5. P. Francois, C. Filsfils, J. Evans, and O. Bonaventure, "Achieving subsecond IGP convergence in large IP networks," *Computer Communication Review*, vol. 35, no. 3, pp. 35–44, 2005.
6. C. Alaettinoglu, V. Jacobson, and H. Yu, "Towards millisecond IGP convergence," November 2000, internet draft, draft-alaettinoglu, ISISconvergence-00.
7. P. Pan, G. Swallow, and A. Atlas, "Fast Reroute Extensions to RSVP-TE for LSP Tunnels," May 2005, Internet RFC 4090.
8. M. Shand and S. Bryant, "IP Fast Reroute Framework," October 2006,
9. A. Shaikh, R. Dube, and A. Varma, "Avoiding Instability during Graceful Shutdown of OSPF," in *Proc. IEEE*.



GLOBAL JOURNALS INC. (US) GUIDELINES HANDBOOK 2011

WWW.GLOBALJOURNALS.ORG

FELLOWS

FELLOW OF INTERNATIONAL CONGRESS OF COMPUTER SCIENCE AND TECHNOLOGY (FICCT)

- FICCT' title will be awarded to the person after approval of Editor-in-Chief and Editorial Board. The title 'FICCT" can be added to name in the following manner e.g. **Dr. Andrew Knoll, Ph.D., FICCT, Er. Pettor Jone, M.E., FICCT**
- FICCT can submit two papers every year for publication without any charges. The paper will be sent to two peer reviewers. The paper will be published after the acceptance of peer reviewers and Editorial Board.
- Free unlimited Web-space will be allotted to 'FICCT 'along with subDomain to contribute and partake in our activities.
- A professional email address will be allotted free with unlimited email space.
- FICCT will be authorized to receive e-Journals - GJCST for the Lifetime.
- FICCT will be exempted from the registration fees of Seminar/Symposium/Conference/Workshop conducted internationally of GJCST (FREE of Charge).
- FICCT will be an Honorable Guest of any gathering hold.

ASSOCIATE OF INTERNATIONAL CONGRESS OF COMPUTER SCIENCE AND TECHNOLOGY (AICCT)

- AICCT title will be awarded to the person/institution after approval of Editor-in-Chief and Editorial Board. The title 'AICCTcan be added to name in the following manner:
eg. **Dr. Thomas Herry, Ph.D., AICCT**
- AICCT can submit one paper every year for publication without any charges. The paper will be sent to two peer reviewers. The paper will be published after the acceptance of peer reviewers and Editorial Board.
- Free 2GB Web-space will be allotted to 'FICCT' along with subDomain to contribute and participate in our activities.
- A professional email address will be allotted with free 1GB email space.
- AICCT will be authorized to receive e-Journal GJCST for lifetime.
- A professional email address will be allotted with free 1GB email space.
- AICHSS will be authorized to receive e-Journal GJHSS for lifetime.

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.



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
<i>Abstract</i>	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
<i>Introduction</i>	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
<i>Methods and Procedures</i>	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
<i>Result</i>	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
<i>Discussion</i>	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
<i>References</i>	Complete and correct format, well organized	Beside the point, Incomplete	Wrong format and structuring

INDEX

A

access · 7, 8, 9, 10, 11, 12, 22, 50, 53, 56, 57
Acoustic · 1, 3
algorithms · 2, 29, 31, 32, 33, 34, 40, 66, 68, 69, 70, 71, 72
analog · 62, 64
applicable · 9, I
approach · 1, 2, 5, 7, 8, 11, 12, 18, 29, 33, 36, 37, 38, 39, 40, 41, 56, 63, 74, 75
Approximation · 66, 71, 72
Architecture · 12, 36, 37, 38, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49, 51, 54, 59
assignments · 8, 11
attribute · 7, 8, 9, 10, 11, 12
authority · 7, 8, 9, 10, 11, 12
authorization · 7, 8, 9, 11, 12

B

Breakdown · 36, 37, 38, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49

C

cascode · 62, 63, 64
certificates · 7, 8, 12
characteristic · 14, 15, 17, 29, 56
circuit · 62, 63, 64, 69
clustering · 29, 31, 33, 34
COCOMO · 37, 38, 47
commercial · 74, 76
communication · 4, 21, 37, 38, 44, 50, 51, 57, 58
Competition · 30
completeness · 66, 69
complexity · 36, 37, 39, 40, 41, 44, 46, 48, 67, 68, 69, 70
Complexity · 39, 40, 46, 71
concentrate · 58
contributions · 42
control · 7, 8, 9, 11, 12, 18, 21, 22, 29, 56, 57, 58
convergence · 31, 74, 75, 76, 77, I
covariance · 1, 3
credentials · 7, 8, 10, 11
criticisms · 45

D

demonstrated · 37, 42, 53, 74

disabilities · 1, 5
distinguishing · 41

E

elaborated · 43
eliminates · 63
establishment · 1, 4, 11, 12, I
extension · 44, 56, 67
extraction · 2, 3, 13, 14, 15, 16, 17, 18, 29, 33

F

framework · 9, 10, 36, 37, 38, 39, 40, 41, 42, 44, 45, 46, 51, 52
Fuzzy · 29, 30, 31, 33, 34, 35, 47

H

Handove · 50
handwritten · 13, 18, 19, 29, 31, 33
Hidden · 1, 2, 3, 4, 5, 6, 12, 13, 14, 15, 16, 17, 18, 19, 20

I

implementation · 7, 8, 9, 11, 18, 29, 36, 38, 39, 41, 52, 57, 58, 62
initializes · 17
integrated · 5, 27, 37, 38, 39, 44, 45, 64
Intelligent · 27, 48, 50, 59

L

Leveraging · 44

M

Markov · 1, 2, 3, 4, 5, 6, 13, 14, 15, 16, 17, 18, 19, 20
mathematical · 13, 14, 18, 50
mechanisms · 8, 56, 75
minimization · 69
Model · 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 33, 37, 47, 50, 58
morphology · 13, 14, 18

N

Network · 13, 21, 22, 23, 24, 25, 26, 27, 28, 50, 51, 52, 53, 54, 56, 57, 58, 71, 74, I
Neural · 13, 14, 15, 16, 17, 18, 19, 20, 29, 33, 34
neurons · 15, 29, 30, 32

O

Organizing · 29, 30, 31, 33, 34, 35
Oriented · 12, 36, 37, 38, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49
orienteering · 70, 71, 72
OUTFC · 74, I

P

parametric · 16, 31
polynomial · 66, 67, 68, 69, 70
protocol · 8, 21, 56, 57, 58, 75, 76, I
provider · 7, 8, 10, 11, 12, 39, 52

R

recognition · 1, 2, 3, 4, 5, 13, 16, 17, 18, 19, 21, 29, 30, 31, 33, 34
resistances · 64
Routing · 54, 66, 67, 68, 69, 70, 71, 72, 73, I

S

Sensor · 21, 22, 23, 24, 25, 26, 27, 28, 56, 57, 58, 59, 60, 61
Simplicity · 40
speech · 1, 2, 3, 5, 16, 19

T

Tifinagh · 13, 14, 15, 16, 17, 18, 19, 20, 29, 30, 31, 33, 34, 35
transmission · 50, 56, 57, 58, 59

U

Ubiquitous · 21, 22, 23, 24, 25, 26, 27, 28
unprecedented · 37

V

vectors · 1, 2, 3, 15, 17, 30, 31
Vehicle · 66, 67, 68, 69, 70, 71, 72, 73
Vertical · 50, 54
voltage · 62, 63, 64

W

wireless · 21, 22, 51, 56, 57, 58



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