

GLOBAL JOURNAL

OF COMPUTER SCIENCE AND TECHNOLOGY: A

Hardware & Computation

Certain Trust Model

Resource Scheduling in Grid

Highlights

A Computational Approach

Enlightenment and Emulation

Discovering Thoughts, Inventing Future

VOLUME 13

ISSUE 2

VERSION 1.0



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY: A
HARDWARE & COMPUTATION

GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY: A
HARDWARE & COMPUTATION

VOLUME 13 ISSUE 2 (VER. 1.0)

OPEN ASSOCIATION OF RESEARCH SOCIETY

© Global Journal of Computer Science and Technology. 2013.

All rights reserved.

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

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

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

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

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

Engage with the contents herein at your own risk.

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

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

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

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

Global Journals Inc.

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

Sponsors: Open Association of Research Society
Open Scientific Standards

Publisher's Headquarters office

Global Journals Headquarters
301st Edgewater Place Suite, 100 Edgewater Dr.-Pl,
Wakefield MASSACHUSETTS, Pin: 01880,
United States of America

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

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

Offset Typesetting

Global Journals Incorporated
2nd, Lansdowne, Lansdowne Rd., Croydon-Surrey,
Pin: CR9 2ER, United Kingdom

Packaging & Continental Dispatching

Global Journals
E-3130 Sudama Nagar, Near Gopur Square,
Indore, M.P., Pin: 452009, 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)

INTEGRATED EDITORIAL BOARD
(COMPUTER SCIENCE, ENGINEERING, MEDICAL, MANAGEMENT, NATURAL
SCIENCE, SOCIAL SCIENCE)

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

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

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

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

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

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

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

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

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

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

Dr. Bart Lambrecht

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

Dr. Carlos García Pont

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

Dr. Fotini Labropulu

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

Dr. Lynn Lim

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

Dr. Mihaly Mezei

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

Dr. Söhnke M. Bartram

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

Dr. Miguel Angel Ariño

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

Philip G. Moscoso

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

Dr. Sanjay Dixit, M.D.

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

Dr. Han-Xiang Deng

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

Dr. Pina C. Sanelli

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

Dr. Roberto Sanchez

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

Dr. Wen-Yih Sun

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

Dr. Michael R. Rudnick

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

Dr. Bassey Benjamin Esu

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

Dr. Aziz M. Barbar, Ph.D.

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

PRESIDENT EDITOR (HON.)

Dr. George Perry, (Neuroscientist)

Dean and Professor, College of Sciences

Denham Harman Research Award (American Aging Association)

ISI Highly Cited Researcher, Iberoamerican Molecular Biology Organization

AAAS Fellow, Correspondent Member of Spanish Royal Academy of Sciences

University of Texas at San Antonio

Postdoctoral Fellow (Department of Cell Biology)

Baylor College of Medicine

Houston, Texas, United States

CHIEF AUTHOR (HON.)

Dr. R.K. Dixit

M.Sc., Ph.D., FICCT

Chief Author, India

Email: authorind@computerresearch.org

DEAN & EDITOR-IN-CHIEF (HON.)

Vivek Dubey(HON.)

MS (Industrial Engineering),

MS (Mechanical Engineering)

University of Wisconsin, FICCT

Editor-in-Chief, USA

editorusa@computerresearch.org

Sangita Dixit

M.Sc., FICCT

Dean & Chancellor (Asia Pacific)

deanind@computerresearch.org

Suyash Dixit

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

President, Web Administration and

Development , CEO at IOSRD

COO at GAOR & OSS

Er. Suyog Dixit

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

SAP Certified Consultant

CEO at IOSRD, GAOR & OSS

Technical Dean, Global Journals Inc. (US)

Website: www.suyogdixit.com

Email: suyog@suyogdixit.com

Pritesh Rajvaidya

(MS) Computer Science Department

California State University

BE (Computer Science), FICCT

Technical Dean, USA

Email: pritesh@computerresearch.org

Luis Galárraga

J!Research Project Leader

Saarbrücken, Germany

CONTENTS OF THE VOLUME

- i. Copyright Notice
 - ii. Editorial Board Members
 - iii. Chief Author and Dean
 - iv. Table of Contents
 - v. From the Chief Editor's Desk
 - vi. Research and Review Papers
-
- 1. E-Commerce Model based on Fuzzy based Certain Trust Model. *1-7*
 - 2. Resource Scheduling in Grid Computing: A Survey. *9-12*
 - 3. A Frame Work for Parallel String Matching- A Computational Approach with Omega Mode. *13-20*
 - 4. A Performance Comparison between Enlightenment and Emulation in Microsoft Hyper-V. *21-28*
-
- vii. Auxiliary Memberships
 - viii. Process of Submission of Research Paper
 - ix. Preferred Author Guidelines
 - x. Index



E-Commerce Model based on Fuzzy Based Certain Trust Model

By Kawser Wazed Nafi, Tonny Shekha Kar, Md. Amjad Hossain
& M. M. A. Hashem

Khulna University of Engineering and Technology, Bangladesh

Abstract- Trustworthiness especially for service oriented system is very important topic now a day in IT field of the whole world. There are many successful E-commerce organizations presently run in the whole world, but E-commerce has not reached its full potential. The main reason behind this is lack of Trust of people in e-commerce. Again, proper models are still absent for calculating trust of different e-commerce organizations. Most of the present trust models are subjective and have failed to account vagueness and ambiguity of different domain. In this paper we have proposed a new fuzzy logic based Certain Trust model which considers these ambiguity and vagueness of different domain. Fuzzy Based Certain Trust Model depends on some certain values given by experts and developers can be applied in a system like cloud computing, internet, website, e-commerce, etc. to ensure trustworthiness of these platforms. In this paper we show, although fuzzy works with uncertainties, proposed model works with some certain values. Some experimental results and validation of the model with linguistics terms are shown at the last part of the paper.

Keywords: *certain trust; fuzzy logic; probabilistic logic; subjective logic; e-commerce trust model.*

GJCST-A Classification : *K.4.4*



Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

E-Commerce Model based on Fuzzy based Certain Trust Model

Kawser Wazed Nafi^α, Tonny Shekha Kar^σ, Md. Amjad Hossain^ρ & M. M. A. Hashem^ω

Abstract- Trustworthiness especially for service oriented system is very important topic now a day in IT field of the whole world. There are many successful E-commerce organizations presently run in the whole world, but E-commerce has not reached its full potential. The main reason behind this is lack of Trust of people in e-commerce. Again, proper models are still absent for calculating trust of different e-commerce organizations. Most of the present trust models are subjective and have failed to account vagueness and ambiguity of different domain. In this paper we have proposed a new fuzzy logic based Certain Trust model which considers these ambiguity and vagueness of different domain. Fuzzy Based Certain Trust Model depends on some certain values given by experts and developers can be applied in a system like cloud computing, internet, website, e-commerce, etc. to ensure trustworthiness of these platforms. In this paper we show, although fuzzy works with uncertainties, proposed model works with some certain values. Some experimental results and validation of the model with linguistics terms are shown at the last part of the paper.

Keywords: *certain trust; fuzzy logic; probabilistic logic; subjective logic; e-commerce trust model.*

1. INTRODUCTION

Trust is a well-known concept in everyday life and often serves as a basis for making decisions in complex situations. There are numerous approaches for modeling trust concept in different research fields of computer science, e.g., virtual organizations, mobile and P2P networks, and E-Commerce. The sociologist Diego Gambetta has provided a definition, which is currently shared or at least adopted by many researchers. According to him "Trust is a particular level of the subjective probability with which an agent assesses that another agent or group of agents will perform a particular action, both before he can monitor such action and in a context in which it affects its own action" [1].

E-commerce is seen as an extension of mail and phone order transactions and is gaining popularity. In E-commerce, business transactions are no longer bound to physical existence, geographic boundaries, time differences or distance barriers. Han and Noh [2] found that several critical failure factors of E-commerce need to

be addressed seriously by the industry to ensure that E-commerce usage will continue to grow. The findings are mainly on the dissatisfaction of customers on the unstable E-commerce systems, a low level of personal data security, inconvenience systems, disappointing purchases, unwillingness to provide personal details and mistrust of the technology. Indeed, customers may doubt the quality of the goods as they may find it difficult to engage in a transaction without proper testing, seeing and touching the products.

To succeed in the fiercely competitive e-commerce marketplace, businesses must become fully aware of Internet security threats, take advantage of the technology that overcomes them, and win customers' trust. Eighty-five percent of Web users surveyed reported that a lack of security made them uncomfortable sending credit card numbers over the Internet. The merchants who can win the confidence of these customers will gain their loyalty—and an enormous opportunity for expanding market share [3].

In person-to-person transactions, security is based on physical cues. Consumers accept the risks of using credit cards in places like department stores because they can see and touch the merchandise and make judgments about the store. On the Internet, without those physical cues, it is much more difficult to assess the safety of a business. Also, serious security threats have emerged. By becoming aware of the risks of Internet-based transactions, businesses can acquire technology solutions that overcome those risks: Spoofing, Unauthorized Disclosure, Unauthorized Action, Eavesdropping and Data Alteration.

In this paper, we are going to propose a newer trust model which is based on fuzzy logic and probabilistic logic. Many trust model developers used "Subjective Logic" for their trust mode. But this subjective model has some problems like: fail to work with uncertain values, lacks of clear mathematical results, etc. Proposed model will solve the problems with "Subjective Logic". In this paper, the mathematical equations are designed such a way that clients and customers can easily understand the output of the model and can take their decision easily, because humans can easily understand fuzzy model's output.

The whole paper is divided in following sections: section II describes the related E-Commerce work of the proposed model, Section III describes the proposed model and its output for different modules, section IV

Authors α σ: Lecturer, Computer Science and Engineering department of Stamford University, Bangladesh. e-mails: kwnafi@yahoo.com, tulip0707051@yahoo.com

Authors ρ ω: Lecturer, Computer Science and Engineering department, Khulna University of Engineering and Technology (KUET), Bangladesh. e-mails: amjad_kuet@yahoo.com, mma.hashem@outlook.com

discusses the experimental results which are worked out in Lab and advantages of the proposed model and section V discusses about the conclusion part of the whole proposed model.

II. RELATED WORK

Several number of ways are there for modeling uncertainty of trust values in the field of trust modeling in Cloud computing and internet based marketing system.[4] But, these models have less capability to derive trustworthiness of a system which are based on knowledge about its components and subsystems. Fuzzy logic was used to provide trust in Cloud computing. Different types of attacks and trust models in service oriented systems, distributed system and so on are designed based on fuzzy logic system [5]. But it models different type of uncertainty known as linguistic uncertainty or fuzziness. In paper [6], a very good model for E-commerce, which is based on fuzzy logic, is presented. But, this model also works with uncertain behavior. Belief theory such as Dempster-Shafer theory was used to provide trust in Cloud computing [7]. But the main drawback of this model is that the parameters for belief, disbelief and uncertainty are dependent on each other. It is possible to model uncertainty using Bayesian probabilities[8] which lead to probability density functions e.g., Beta probability density function. It is also possible to apply the propositional standard operators to probability density functions. But this leads to complex mathematical operations and multi-dimensional distributions which are also hard to interpret and to visualize. An enhanced model recently being developed for using in Cloud computing is known as Certain Trust. This model evaluates propositional logic terms that are subject to uncertainty using the propositional logic operators AND, OR and NOT[1].

Manchala [9] proposes a model for the measurement of trust variables and the fuzzy verification of E-Commerce transactions. He highlights the fact that trust can be determined by evaluating the factors that influence it, namely risk. He defines cost of transaction, transaction history, customer loyalty, indemnity and spending patterns as the trust variables. But he fails to solve the following problems of E-Commerce: - Suitable variables for outputs, establish relations between variables and fails to support theoretical logics for E-Commerce trust models. Jøsang[10] also works with trust models and work with "Subjective Logic", but this model is fully depended with uncertainty. S. Nefti proposed a model which solves the problem of Manchala and Jøsang, but fails to solves problems related with uncertainty and can't show behavioral probability of a E-Commerce based company.

III. PROPOSED TRUST MODEL FOR E-COMMERCE

Certain Trust Model was constructed for modeling those probabilities, which are subject to uncertainty. This model was designed with a goal of evidence based trust model. Moreover, it has a graphical, intuitively interpretable interface [1] which helps the users to understand the model (Figure 1). The representational model focuses on two crucial issues:

- How trust can be derived from evidence considering context-dependent parameters?
- How trust can be represented to software agents and human users?

For the first one, a relationship between trust and evidence is needed. For this, they had chosen a Bayesian approach. It is because it provides means for deriving a subjective probability from collected evidence and prior information [1]. At developing a representation of trust, it is necessary to consider to whom trust is represented. It is easy for a software component or a software agent to handle mathematical representations of trust. For it, Bayesian representation of trust is appropriate. The computational model of Certain Trust proposes a new approach for aggregating direct evidence and recommendations. In general, recommendations are collected to increase the amount of information available about the candidates in order to improve the estimate of their trustworthiness. This recommendation system needs to be integrated carefully for the candidates and for the users and owner of cloud servers. This is called robust integration of recommendations. In order to improve the estimate of the trustworthiness of the candidates, it is needed to develop recommendation system carefully. This is called robust integration of recommendations [1].

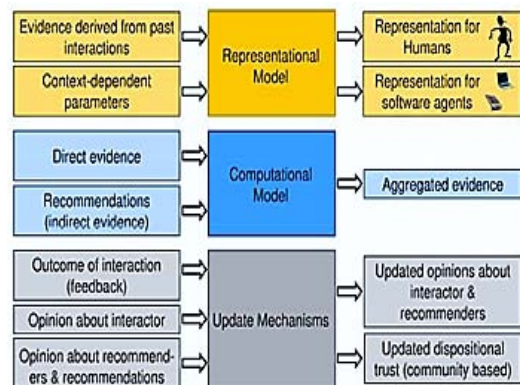


Figure 1 : Block Diagram of Trust models

Three parameters used in certain logic: average rating t , certainty c , initial expectation f . The average rating t indicates the degree to which past observations support the truth of the proposition. The certainty c indicates the degree to which the average rating is assumed to be representative for the future. The initial

expectation f expresses the assumption about the truth of a proposition in absence of evidence [1].
The equations for these parameters are given below:-

$$\text{Equation for average rating } t = \begin{cases} 0.5 & \text{if } r + s = 0 \\ \frac{r}{r+s} & \text{else} \end{cases} \quad (1)$$

Here, r represents number of positive evidence and s represents number of negative evidence defined by the users or third person review system.

$$\text{Equation for certainty, } C = \frac{N.(r+s)}{2.w.(N-(r+s))+N.(r+s)} \quad (2)$$

Here, w represents dispositional trust which influences how quickly the final trust value of an entity shifts from base trust value to the relative frequency of positive outcomes and N represents the maximum number of evidence for modeling trust. Using these parameters the expectation value of an opinion $E(t, c, f)$ can be defined as follows:

$$E(t, c, f) = t * c + (1 - c) * f \quad (3)$$

The parameters for an opinion $o = (t, c, f)$ can be assessed in the following two ways: direct access and Indirect access. Certain Trust evaluates the logical operators of propositional logic that is AND, OR and NOT. In this model these operators are defined in a way that they are compliant with the evaluation of propositional logic terms in the standard probabilistic approach. However, when combining opinions, those operators will especially take care of the uncertainty that is assigned to its input parameters and reflect this (un)certainly in the result. The definitions of the operators as defined in the CTM are given in the Table I.

Table 1: Definition of Operators

OR	$C_{A \vee B} = C_A + C_B - C_A C_B - \frac{c_A(1-c_B)f_B(1-t_A) + (1-c_A)c_B f_A(1-t_B)}{f_A + f_B - f_A f_B}$
	$t_{A \vee B} = \begin{cases} \frac{1}{C_{A \vee B}} (C_A t_A + C_B t_B - C_A C_B t_A t_B) & \text{if } C_{A \vee B} \neq 0 \\ 0.5 & \text{else} \end{cases}$
	$f_{A \vee B} = f_A + f_B - f_A f_B$
AND	$C_{A \wedge B} = C_A + C_B - C_A C_B - \frac{(1-c_A)c_B(1-f_A)t_B + c_A(1-c_B)(1-f_B)t_A}{1-f_A f_B}$
	$t_{A \wedge B} = \begin{cases} \frac{1}{C_{A \wedge B}} (C_A C_B t_A t_B + \frac{c_A(1-c_B)(1-f_A)f_B t_A + (1-c_A)c_B f_A(1-f_B)t_B}{1-f_A f_B}) & \text{if } C_{A \wedge B} \neq 0 \\ 0.5 & \text{else} \end{cases}$
	$f_{A \wedge B} = f_A f_B$
NOT	$t \neg_A = 1 - t_A, \quad c \neg_A = 1 - c_A \quad \text{and} \quad f \neg_A = 1 - f_A$

With the help of these parameters and operators derived from Certain Trust, we have defined two new parameters, Trust T and behavioral probability, P [11]. Trust T is calculated from certainty c and average rating t . the equation is:

$$\text{Trust, } T = \frac{c * t}{\text{High scaling value of rating}} * 100\% \quad (4)$$

Here, High scaling value of rating means the upper value of the range of rating.

Calculating T , we have applied FAM rule of fuzzy logic for creating a relation between certainty c and average rating t . Trust T represents this relation in percentage such a way that the quality of the product can easily be understood. Another parameter, behavioral probability, P , represents how the present behavior of the system varies from its initially expected value and it is proposed with the help of probabilistic logic. It may be less, equal or higher than the initial expectation given by the system developer or the manager of the office. The equation for P is:

$$\text{Behavioral probability, } P = \frac{(T)-f}{f} * 100\% \quad (5)$$

If $T < f$, lower probability to show expected behavior.

If $T > f$, higher probability to show expected behavior.

If $T = f$, balanced with the expected behavior.

According to Gaussian, the membership function for fuzzy input sets depends on two types of parameter, standard deviation σ and mean c . The equation for membership function is:-

$$f(x; \sigma; c) = \exp\left(\frac{-(x-c)^2}{2\sigma^2}\right) \quad (6)$$

Fuzzy inference is the process of formulating the mapping from a given input to an output using fuzzy logic. The mapping then provides a basis from which decisions can be made, or patterns discerned. There are two concepts of fuzzy logic systems. They are: - linguistic variables and fuzzy if then else rule. The linguistic variables' values are words and sentences where if then else rule has two parts; antecedents and consequent parts which contain propositions of linguistic variables. Numerical values of inputs $x_i \in U_i (i = 1, 2, \dots, n)$ are fuzzified into linguistic values, $F_1, F_2, \dots, F_n$. Here F_i denotes the universe of discourse $U = U_1 * U_2 * \dots * U_n$. The output linguistic variables are $G_1, G_2, \dots, G_n$. The if-then-else rule can be defined as: -

$$R^0: \text{IF } x_i \in F_i^j \text{ and } \dots \text{ and } x_n \in F_n^j \text{ THEN } y \in G_j. \quad (7)$$

The trust relationships among customers and vendors are hard to assess due to the uncertainties involved. Two advantages of using fuzzy-logic to quantify trust in E-commerce applications are: (1) Fuzzy inference is capable of quantifying imprecise data and quantifying uncertainty in measuring the trust index of the vendors. (2). Fuzzy inference can deal with variable dependencies in the system by decoupling dependable variables. The general trust model proposed by s.nemfi [12] is composed of five modules. Four modules will be used to quantify the trust measure of the four factors identified in our trust model (Existence, Affiliation, Policy and Fulfilment). The fifth module will be the final decision maker. Figure II describes general trust model for E-commerce.

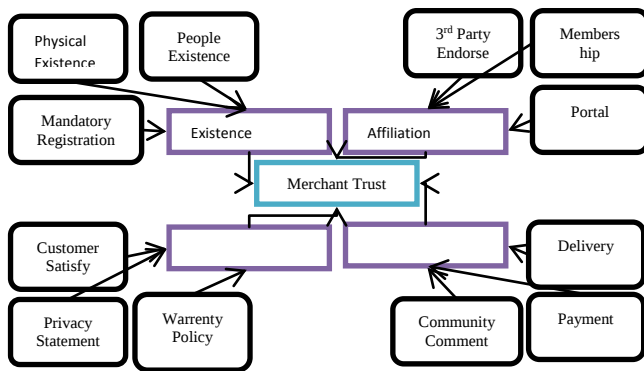


Figure 1 : Trust Model for E-commerce

The inputs of the first module are Physical Existence, People Existence and Mandatory registration and the output of the affiliation module. For module 2, the inputs are Third Party Endorsement, Membership and Portal and output is affiliation variables. Module 3 inputs are Customer Satisfaction, Privacy and Warranty for Policy. Finally, the fourth module inputs are represented by Delivery, Payment Methods and Community Comment. The Final module combines the four modules which are Existence, Fulfillment, Policy and the Affiliation for producing the final output Merchant Trust. Trust values of every element of each module are calculated with the help of parameter c and t and then average values of them work as output of each model. We divide the classes of all input and output variables into five parts according to acceptable degrees of survey.

Fuzzy specification rules, according to equation 8, used for Existence module are given below:-

1. If (People_Eistence is Very_Low) and (Physical_Existence is Very_Low) and (Mendatory_Registration is Very_Low) then (Existence is Very_Low) (1).
2. If (People_Eistence is Low) and (Physical_Existence is Very_Low) and (Mendatory_Registration is Very_Low) then (Existence is Very_Low) (1).
-
- 130.If (People_Eistence is Very_High) and (Physical_Existence is Very_High) and (Mendatory_Registration is Very_High) then (Existence is Very_High) (1).

With the help of fuzzy linguistic variables, we get the final output of Existence modules given in Figure 3.

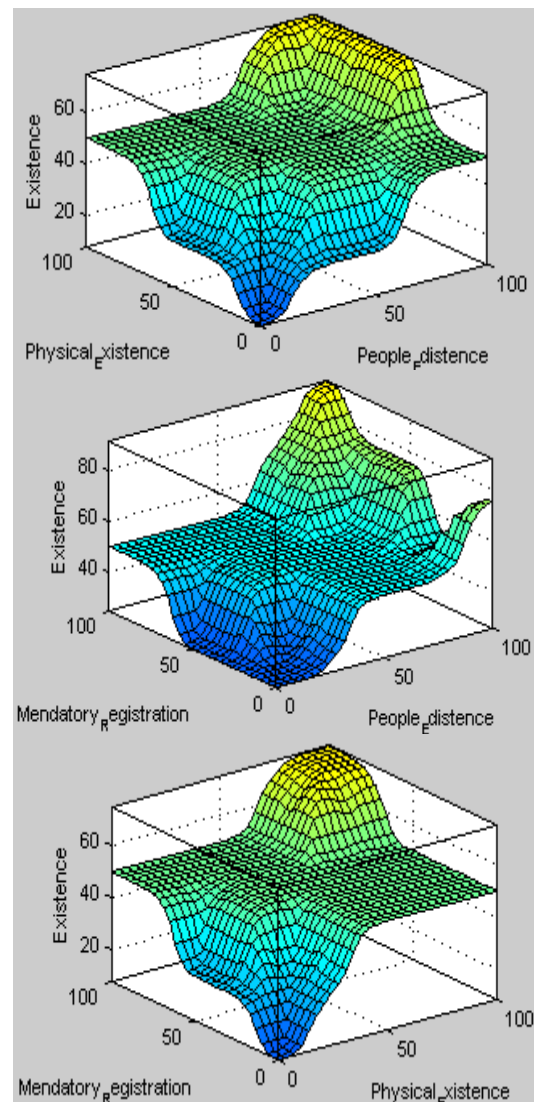


Figure 3 : Mapping Surface for Existence Module

After getting fuzzy outputs from 4 modules of trust model, we have used them as input for our final module, Merchant trust and have got the membership functions shown in Figure 5 and outputs for different variables shown in Figure 4.

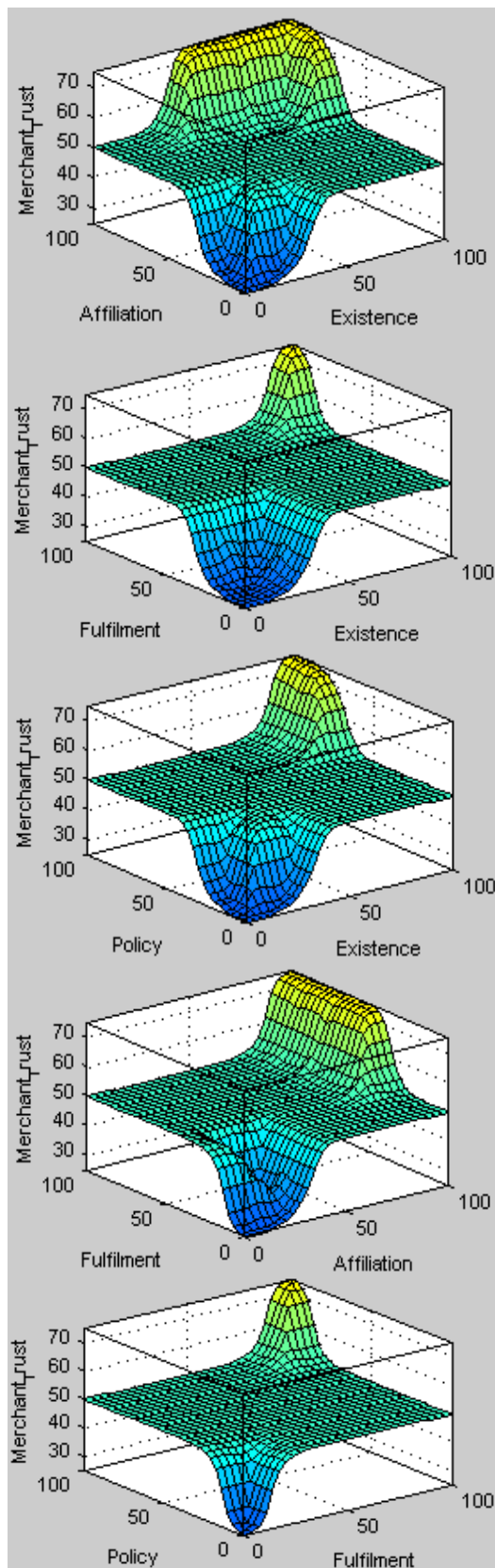


Figure 4 : Mapping Surface for Merchant Trust Module

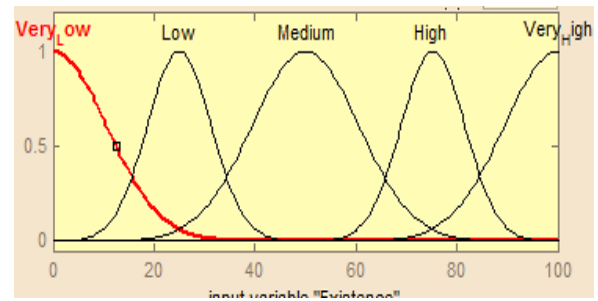


Figure 5 : Membership Functions for Final Output Merchant Trust

IV. EXPERIMENTAL RESULT

In this section we are going to show the effects of Fuzzy Based Certain Trust model in the field of E-Commerce Architecture. We have run this model in our lab for several times and have got different types of trust values for different environments. We are going to show two scenarios consisting 100 people each and trying to show the evaluating process of our proposed model.

a) Case Study 1

Let us consider a company, name A, running an online business following our E-Commerce structure. According to Certain Trust Model's operators defined in equation 3, 4, and 5, we know that, the input for this model is r , s , f and w . Let, for any time the input values are $r=5$, $s=2$, $f=0.5$ and $w=1$ where, number of evidences are $N=7$. Then, the output values are:- average rating $t = 0.714$ and $c=0.724$ and then, $E = 0.65$. Now, for mapping it to our proposed model, we need to modify t . Here,

$$t' = t * \text{scale of rating} \quad (8)$$

Usually, the scale of rating is 5. Now, the new average rating is $t = 0.714 * 5 = 3.57$. Then, the value of parameter Trust, $T = ((3.57 * 0.724) / 5) * 100 = 51.69\%$. Let, considering a customer has come to buy some products from Company A website. After completing all his prerequisites for completing a full transaction with this company, he needs to give his recommendation for different steps of 4 modules. Again, considering the previous states of Certain Trust Model, the present condition of Trust Values of the company are given in Table II.

Table 2 : Trust Values of Different Modules

Module No	Variables	Values	Trust Values	Final Trust
Module 1 Existences	Physical Existence	$c=0.6$, $t=3.5$	42%	43%
	People Existence	$c=0.3$, $t=4$	24%	
	Mandatory Existence	$c=0.7$, $t=4.5$	63%	
Module 2	Third Party Endorsement	$c=0.6$, $t=4.25$	51%	

Affiliation	Membership	c=0.9, t=4.8	86.4%	70%
	Portal	c=0.8, t=4.5	72%	
Module 3 Fulfillment	Delivery	c=0.5, t=4.35	43.5%	59.5%
	Payment	c=0.8, t=4.30	69%	
	Community Customer	c=0.7, t=4.7	66%	
Module 4 Policy	Customer Satisfaction Policy	c=0.8, t=4.6	73.6%	61%
	Privacy Statement	c=0.7, t=4.5	63%	
	Warranty Policy	c=0.5, t=4.6	46%	

Now, the Trust Values for Merchant Trust for this transaction is $T = ((43+70+59.5+61)/4) = 58.375\%$. So, considering other recommendation systems' trust values and previous trust values, the Trust of the company will be clustered in medium trust values class.

b) Case Studies 2

Let us again consider a newer company, name B. This company uses the proposed model and tries to calculate their trust values for representing their Trustworthy condition to clients of the whole world. Considering all values and equations, they get the values shown in Table III.

Table 3 : Trust Values of Different Modules

Module No.	Variables	Values	Trust Values	Final Trust
Module 1 Existences	Physical Existence	c=0.5, t=3.8	38%	57%
	People Existence	c=0.76, t=4.25	64.6%	
	Mandatory Existence	c=0.72, t=4.75	68.4%	
Module 2 Affiliation	Third Party Endorsement	c=0.5, t=4.5	45%	39%
	Membership	c=0.9, t=4.8	86.4%	
	Portal	c=0.75, t=4.25	63.75%	
Module 3 Fulfillment	Delivery	c=0.46, t=4.45	41%	50.5 %
	Payment	c=0.6, t=4.7	56.4%	
	Community Customer	c=0.56, t=4.86	54%	
Module 4 Policy	Customer Satisfaction Policy	c=0.42, t=4.8	40%	42.55 %
	Privacy Statement	c=0.65, t=3.95	51.65%	
	Warranty Policy	c=0.55, t=3.26	36%	

Now, the Trust Values for Merchant Trust for this transaction is $T = ((57+39+50.5+42.55)/4) = 47.26\%$.

Now, Behavioral probability of company A is $P = (\frac{58-5}{.5}) * 100\% = 16\%$ Up than base. Considering Base $f=1.5$. And Behavioral Probability of company B is $P = (\frac{47-5}{.5}) * 100\% = 6\%$ lower than Base. Now, the newer client can easily distinguish between these two companies and can take decision easily with which company he will start his deal.

The advantages of proposed model over S.Nefti are given below:

Table 4 : Advantages of Proposed Model

Points of Discussion	S.Nefti Model	Proposed Model
Number of Fuzzy Classes	2-3 Membership function, So, lower specification	5 Membership functions, So, classification and specification is Higher
Dependency	Depends on Uncertain Fuzzy System	Depends on Certain Fuzzy System
Behavioral Probability	Not Present	Present, so present condition can easily understandable
Mapping in Cloud Architecture	Not Applicable	Fully Applicable

V. CONCLUSION

In this paper, we presented a system based on fuzzy logic based certain trust model to support the evaluation and the quantification of trust in E-commerce. As stated in many trust models, there are other aspects that contribute to the completion of online transactions. This includes the price, the rarity of the item and the experience of the customer. In this paper, we mainly concentrate on the structure of the E-commerce companies and different elements of it. With the help of proposed model, one can easily distinguish between companies and other elements of trust model will also be considered by the clients rating for the companies and for different transactions.

We have some new idea to imply in our proposed model in future. Firstly, we want to apply evolutionary algorithm with this model to optimize and design the rules. We want to apply price comparison as a parameter for a product in our model for ensuring accurate trust measuring model for a normal e-commerce website. Secondly, more development of behavioral probability parameter so that it can directly prohibit different types of security breaking questions like Sybil attack, false rating, etc.. At present, this parameter works indirectly with security options. Last of all, we want to work with all the points of E-Commerce architecture in

our future model so that only our proposed model can fulfill all the requirements of customers who rely on online transactions for their business.

VI. ACKNOWLEDGMENT

The authors wish to thank Sebastian Ries, Sheikh MahbubHabib, Max Mühlhäuser and Vijay Varadharajan for their renowned work and thinking. Again, want to thank all the researchers and developers of Trust models for their excellent work which make authors work easier to accomplish.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Ries, S.: Trust in Ubiquitous Computing. PhD thesis, Technische University at Darmstadt, pp: 1-192, 2009.
2. K.S. Han and M.H. Noh, "Critical Failure Factors that Discourage the Growth of Electronic Commerce", *International Journal of Electronic Commerce*. 4(2), pp. 25-43, 1999.
3. "Building an E-Commerce Trust Infrastructure SSL Server Certificates and Online Payment Services", *Verisign*, pg: 1-37.
4. Mohammed Alhamad, Tharam Dillon, and Elizabeth Chang A Trust-Evaluation Metric for Cloud applications *International Journal of Machine Learning and Computing*, Vol. 1, No. 4, October 2011.
5. Jøsang, A., Ismail, R., Boyd, C.: A survey of trust and reputation systems for online service Provision. *Decision Support Systems* 43(2) (2007) 618–644.
6. Vojislav Kecman, "Learning and Soft computing", Google online library, 2001.
7. Kerr, R., Cohen, R.: Smart cheaters do prosper: defeating trust and reputation systems. In: AAMAS '09: Proceedings of The 8th International Conference on Autonomous Agents and Multiagent Systems, (2009) 993–1000.
8. L. Zadeh. "Fuzzy sets", *Journal of Information and Control*, 8:338—353, 1965.
9. D.W. Manchala, "E-Commerce Trust Metrics and Models", *IEEE Internet Computing*, March-April 2000, pp. 36-44.
10. A. Jøsang, "Modelling Trust in Information Society", Ph.D. Thesis, Department of Telematics, Norwegian University of Science and Technology, Trondheim, Norway, 1998.
11. Kawser Wazed Nafi, Tonny Shekha Kar, Md. Amjad Hossain, M. M. A. Hashem, "An Advanced Certain Trust Model Using Fuzzy Logic and Probabilistic Logic theory", *IJACSA*, vol 3, no 11, 2012.
12. Samia, Nefti, FaridMeziane, KhairudinKasiran, A Fuzzy Trust Model for E-Commerce, 2004.

This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 2 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Resource Scheduling in Grid Computing : A Survey

By Sonal Nagariya & Mahindra Mishra

LNCT/RGPV, India

Abstract- Grid computing is a computing framework to meet growing demands for running heterogeneous grid enabled applications. A grid system is composed of computers which are separately located and connected with each other through a network. Grids are systems that involve resource sharing and problem solving in heterogeneous dynamic grid environments. Here we present five different approaches/algorithms for resource allocation/ Scheduling in grid computing environment.

Keywords: *grid computing, dynamic environment, resource allocation, genetic algorithm, artificial neural networks, priority, predictions.*

GJCST-A Classification : *C.2.4*



Strictly as per the compliance and regulations of:



Resource Scheduling in Grid Computing: A Survey

Sonal Nagariya^α & Mahindra Mishra^σ

Abstract- Grid computing is a computing framework to meet growing demands for running heterogeneous grid enables applications. A grid system is composed of computers which are separately located and connected with each other through a network. Grids are systems that involve resource sharing and problem solving in heterogeneous dynamic grid environments. Here we present five different approaches/algorithms for resource allocation/ Scheduling in grid computing environment.

Keywords: grid computing, dynamic environment, resource allocation, genetic algorithm, artificial neural networks, priority, predictions.

I. INTRODUCTION

Increased network bandwidth, more powerful computers, and the acceptance of the Internet have driven the ongoing demand for new and better ways to compute. Organizations & institutions alike continue to take advantage of these advancements, and constantly seek new technologies and practices that enable them to reinvent the way they conduct business. [1]. The Grid environment is the IT infrastructure of the future promises to transform computation, communication, and collaboration (see figure 1). Depending on the main type of service a grid offers, it can be classified as: computational grid, access grid, data grid or data centric grid [2].

Computational grids provide high performance computing; access grids allow the access to a specific resources; data grids store and move large data sets; and data centric grids enable distributed repositories of data that cannot be stored in a single one. Grid computing is like a distributed computing applies the resources of many computers in a network to solve single problem at the same time-usually to a scientific or technical problem that requires a great number of computer processing cycles or access to large amounts of the efficiency of underlying system by maximizing utilization of distributed heterogeneous resources.

Resource allocation is one of the important system services that have to be available to achieve the objectives from a grid computing. A common problem arising in grid computing is to select the most efficient resource to run a particular program. Also users are

requiring reserving in advance the resources needed to run their programs on the grid [3]. In this paper we review the various Dynamic resource allocation algorithms like ELM, ANN, Priority, performance prediction, Genetic algorithm for Grid Computing Environment.

This paper is organized as follows: Section II introduces the problem of resource allocation in grid networks. Section III presents a survey of different resource allocation algorithms in current era. Section IV conclusion.

II. THE PROBLEM OF RESOURCE ALLOCATION IN GRID ENVIRONMENT

Before scheduling the resources in the grid environment, the characteristics of the grid should be taken into account. Some of the characteristics of the grid include:

1. Geographical distribution of the resources
2. Heterogeneity, a grid consists of hardware as well as software resources that may be files, software components, sensor programs,
3. Resource sharing, between different organizations
4. Multiple administrations, where each organization may establish their own different security and administrative policies to access their resources
5. Resource coordination between heterogeneous computing power
6. Scheduling is highly complicated by the distributed ownership of the grid resources as Load balancing.

The problem under such resource allocation algorithms are-

1. An estimation of the computational load of each job.
2. The computing capacity of each resource.
3. An estimation of the prior load of each one of the resources is required.
4. An estimation on maximum completion time(MTC).

On the basis of these key points we survey different types of dynamic resource scheduling algorithms.

III. SURVEY OF DIFFERENT RESOURCE SCHEDULING ALGORITHM

This section provides a brief overview of five existing resource scheduling Algorithms is pursued in grid environment.

Authors α σ: Department of Information Technology, Lakshmi Narain college of Technology, R.G.P.V., Bhopal.
e-mails: sonalnagariya@yahoo.com,
mahendra.mishra02@gmail.com

a) ELM

In a Grid environment, there are two levels of resource allocation, one for Grid schedulers and the

other for local resource schedulers [4]. Before selecting a resource, Grid scheduler does a minimum requirement filtering.

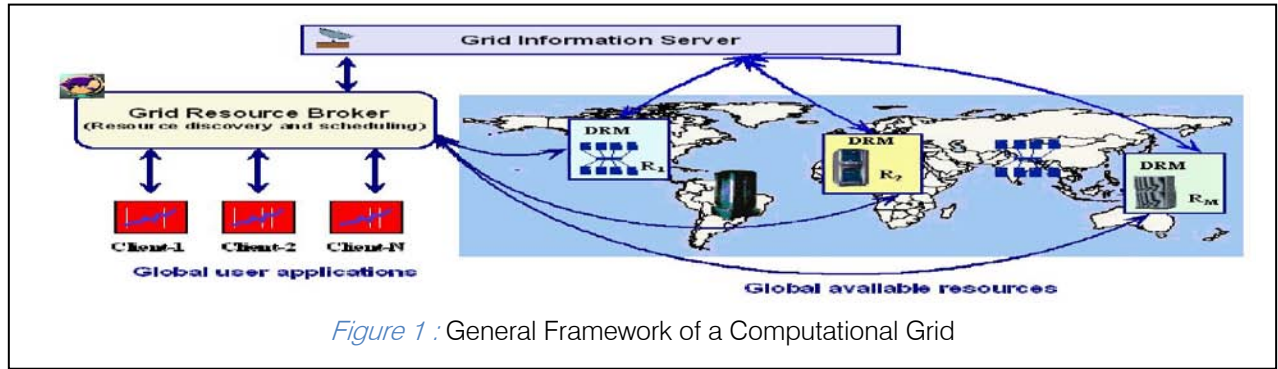


Figure 1 : General Framework of a Computational Grid

which shortlists available resources by comparing resource requirement and job property.

$$\Phi(\text{preference}): R \times J \rightarrow N1 \quad (1)$$

Where $J = \{J\}$ referring to the set consisting of the job to be scheduled. Note that $|J| = 1$, here resource allocation is considered only one job at a time [5].

But there is a disadvantage is that this policy requires users to have sufficient prior knowledge for a particular problem in order to determine the ranking function.

But use of ELM (fast learning algorithm for single-hidden layer feed forward neural networks (SLFNs))(see figure 2); with prediction resource allocation policy will inherit the feature of fast learning and perform satisfactory.

A prediction-based resource allocation policy is defined as,

$$\Phi: \hat{k} = \arg \max P(R_k, J), \text{ for } 1 \leq k \leq |R| \quad (2)$$

where $P: R \times J \rightarrow R$ denotes a predicting function which maps resource and job to a performance value for the underlying preference. In this method machine learning lifts the requirement of sufficient prior knowledge from users and learning can provide accurate and timely results. The training of ELM is currently based on supervised learning, and thus training data needs to be labeled. However, for real-world Grid applications labeled data may not be available all the time. In this approach there is a filter with minimum requirements, if process doesn't full-fill the minimum requirements will be rejected.

b) ANN

Artificial neural networks (ANN) are attempts to reproduce human brain potentialities in a small scale, specially its learning ability. In this approach neuron mathematical model proposed by McCulloch and Pitts (Multi layer perceptron) are used. McCulloch and Pitts had a binary output and inputs, each of different

excitatory or inhibitory gains, known as Synaptic weights (or weights) [6]. The values of input signals and related weights determined neuron output. Through this model automatically capture the requirements of the user and use it for resource selection (see figure 3).

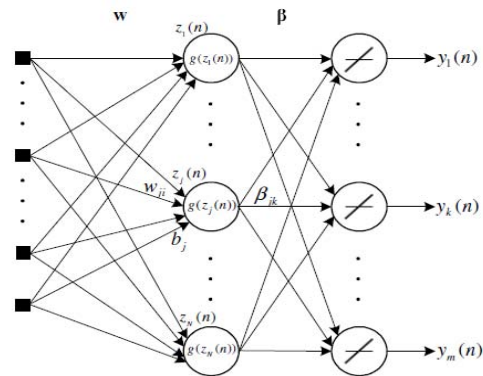


Figure 2 : A sketch of ELM neural network

The effectiveness of the scheduling methods assessed and evaluated using evaluation metrics like make span and flow time. Make span is the time taken by the grid system to complete the latest task; and flow time is the sum of execution times for all the tasks presented to the grid [7-8]. But there is a disadvantage is that this policy requires users to have sufficient prior knowledge for a particular problem. It reduces either make span or flow time.

IV. PRIORITY SCHEDULER

Priority scheduler completed a task by using highly utilized low cost resources with minimum computational time. Our scheduling algorithm uses the priority queue of resources to achieve a higher throughput (See figure 4). This algorithm is performing better for task in real environment. In this algorithm initially setting the load factor of each resource is zero and priority of each resource is one because there is no process for execution. As the process arrives the load

factor of resources increased and priority of resources [9].

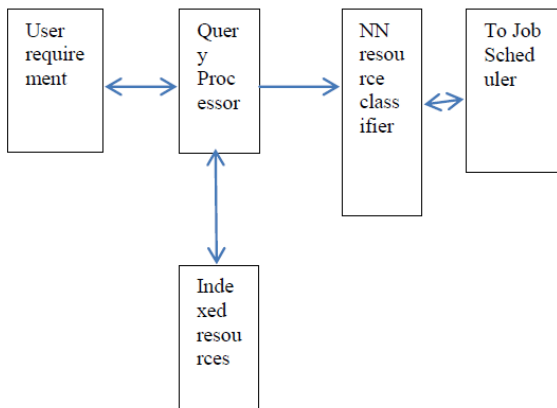


Figure 3 : Grid resource selector

Let $T(P_i, R_j)$ be the total cost for i^{th} process in j^{th} resources can be calculated as-

$$\sum_{i=0}^m \sum_{j=0}^n T(P_i, R_j) = \sum_{i=0}^m \sum_{j=0}^n t_i \times PR + CT \quad (3)$$

Where t_i is the execution time of process, PR =Priority Number, CT =Communication Time

But still grid application performance remains a challenge in dynamic grid environment because resources can be submitted to Grid and can be withdrawn from Grid at any moment. This characteristic of Grid makes Load Balancing one of the critical features of Grid infrastructure there is another drawback is it is basically calculate load factor on there sources through the total no of jobs but it not concern their properties or related information.



Figure 4 : Shows priority VS. Load factor

V. PERFORMANCE PREDICTION

Performance prediction is an Analytical performance model. This system is used with a gravitational wave physics experiment, LIGO(Laser Interferometer gravitational wave observatory) for which initial results indicate an average age of 24% reduction in execution time using the performance prediction versus a random selection of resources [10]. In this method Fourier transformation with prophesy function is used .Performance is determined by interaction between the dataflow and computations in the application. (See figure 5) The Prophesy uses the predicted execution times to rank each candidate site. The less execution time is, the higher the site is ranked. The probabilities

are also based on the predicted Execution times. Assume that the list of candidate sites is $S_1, S_2, \dots, S_n$ and their predicted execution times are $T_1, T_2, \dots, T_n$ respectively[10]. We use the following equation to calculate the selection probability of the site S_i :

$$\frac{1}{T_i} = \frac{1}{\sum_{j=1}^n T_j} \quad (4)$$

The main drawback is that, it holds total number of resources till the process is completed. Another drawback is that it is basically works on historical data so reusing the existing data product is not always the best choice in terms of the total execution time and it is non-adaptive in nature.

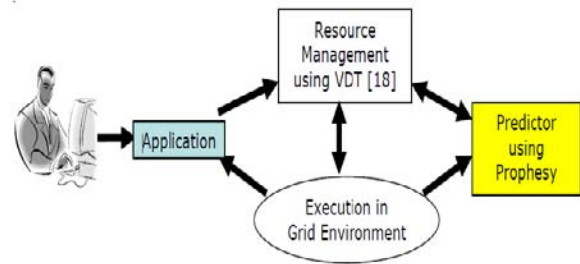


Figure 5 : System workflow

VI. GENETIC ALGORITHM

Heuristic approach like genetic algorithm is used to solving the resource allocation in the grid. Genetic algorithm is used to solve the both the unit commitment problem and the heterogeneous resource Scheduling problem for grid computing [12]. Resource Allocation problem is represented in a chromosome. The chromosome is made up of genes, each representing one asset of the system. Each solution has a fitness value associated with it. These values are used to evaluate the chromosomes that are trying to be optimized. Group of chromosome in the population may violate a constraint because it might be able to produce a child that performs well and fixes the constraint violation as it evolves. There are many different ways to perform crossover, selection and mutation to generate new Childs. (See figure 6) As the number of generations increases, however, the genetic algorithm is less likely to keep a chromosome that violates any of the constraints. In the initial population, it is usually beneficial to have some type of genetic preconditioning in the form of seeding. This seeding uses some number of solutions in the initial population that are not generated at random. Initial seeds that do not violate any constraints. The main benefit of genetic algorithm is that it is adaptive in nature, and reduce make span and flow time both .The fundamental criterion is that of minimizing the make span, that is, the time when finishes the latest job. A secondary criterion is to minimize the flow time that is,

minimizing the sum of finalization times of all the jobs. These two criteria are defined as follows:

$$\begin{aligned} \text{Make span: } \min_{s \in \text{Sched}} \{ \max_{j \in \text{Jobs}} F_j \} \text{ and} \\ \text{Flow time: } \min_{s \in \text{Sched}} \{ \sum_{j \in \text{Jobs}} F_j \} \end{aligned} \quad (5)$$

Where F_j denotes the time when job j finalizes, Sched is the set of all possible schedules and Jobs the set of all jobs to be scheduled [11]. The drawback of the genetic algorithm is that it provides optimal solution but doesn't provide the best solution.

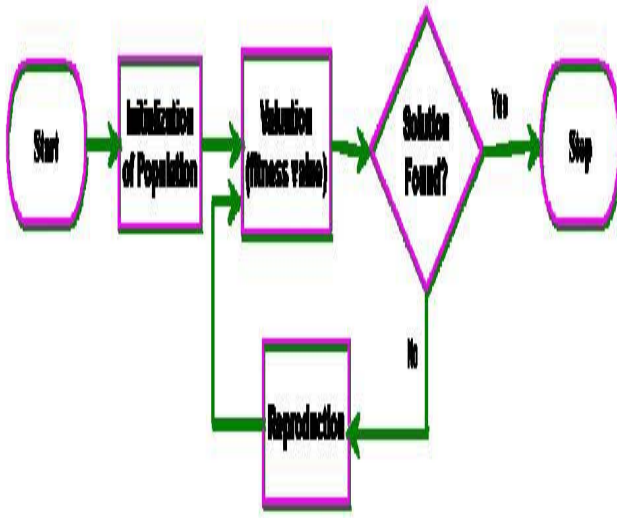


Figure 6 : Flowchart of GA Iteration

VII. CONCLUSION

In heterogeneous grid environments, availability of resources may possibly fluctuate so resource allocation is an NP complete problem where there is no final solution. In this review we have find different approach about resource allocation. Based on dynamic resources allocation technique. Our main objective is to review the various grid resource scheduling strategies which will in turn serve as a guide for researchers We seen that Priority allocation of grid computational are not successful method for resource utilization. For the ELM ANN and Performance Prediction provide best solution but real-world Grid applications labeled data may not be available all the time. For the application of meta-heuristic function improve the performance of the time span of process execution and determine the failure of process. The genetic algorithm works on both flows time and make span, but it provide optimal solution .we can be analyzed for problem of resource allocation such that new research could be focused to produce better solution by improving the effectiveness and reducing the limitations. In future our focus on developing an efficient resource allocation algorithm which not only reduce the communication time of applications with adaptability of resources in Grid computing environment.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Akash K Patel, Kinjal A Faldu , Meghna R Goswami , Mehta Prashant, "Grid Computing: An Overview", Volume 3, Issue 2, February 2013) pp-602.
2. D. B. Skillicorn, "Motivating Computational Grids," in 2nd IEEE/ACM International Symposium on Cluster Computing and the Grid (CCGRID'02), May 2002, pp. 401–406. ISSN: 2277-3754 ISO 9001:2008 Certified
3. Darshan Kanzariya, Sanjay Patel, "Survey on Resource Allocation in Grid", International Journal of Engineering and Innovative Technology (IJEIT) Volume 2, Issue 8, February 2013.
4. J. M. Schopf., "Ten actions when grid scheduling", *International Series in Operations Research and Management Science*, pages 15–24, 2003.
5. Guopeng Zhao, ZhiqiShen, Chunyan Miao, "A Fast and Intelligent Resource Allocation Service for Service-Oriented Grid", 15th International Conference on Parallel and Distributed Systems.
6. T. R. Srinivasan, R. Shanmugalakshmi, "Neural Approach for Resource Selection with PSO for Grid Scheduling", *International Journal of Computer Applications* (0975 – 8887) Volume 53–37 No.11, September 2012.
7. Izakian, H., et al., "A novel particle swarm optimization approach for grid job scheduling." *Information Systems, Technology and Management*, 2009: p. 100-109.
8. Sayadi, M.K., R. Ramezani, and N. Ghaffari-Nasab, "A discrete firefly meta-heuristic with local search for makespan minimization in permutation flow shop scheduling problems." *International Journal of Industrial Engineering Computations*, 2010. 1: p. 1-10.
9. Mayank Kumar Maheshwari, Abhay Bansal, "Process Resource Allocation in Grid Computing using Priority Scheduler", *International Journal of Computer Applications* (0975 – 8887) Volume 46– No.11, May 2012, 20.
10. Seung-Hye Jang, Xingfu Wu, Valerie Taylor, Gaurang Mehta, Karan Vahi, EwaDeelman, "Using Performance Prediction to Allocate Grid Resources", *GriPhyN Technical Report* 2004-25.
11. Javier Carretero, FatosXhafa, Ajith Abraham, "GENETIC ALGORITHM BASED SCHEDULERS FOR GRID COMPUTING SYSTEMS", *International Journal of Innovative Computing, Information and Control* ICIC International 2005 ISSN 1349-4198 Volume 3, Number 6, December 2007 pp. 0–0.
12. Tim Hansen, Robin Roche, Siddharth Suryanarayanan, Howard Jay Siegel, Daniel Zimmerle, Peter M. Young, Anthony A. Maciejewski, "A Proposed Framework for Heuristic Approaches to Resource Allocation in the Emerging Smart Grid".



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 2 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Frame work for Parallel string Matching- A Computational Approach with Omega Model

By K Butchi Raju, Chinta Someswara Rao, Dr. S. Viswanadha Raju

GRIET, India

Abstract- Now a day's parallel string matching problem is attracted by so many researchers because of the importance in information retrieval systems. While it is very easily stated and many of the simple algorithms perform very well in practice, numerous works have been published on the subject and research is still very active. In this paper we propose a omega parallel computing model for parallel string matching. Experimental results show that, on a multi-processor system, the omega model implementation of the proposed parallel string matching algorithm can reduce string matching time by more than 40%.

Keywords: *string matching; parallel string matching; computing model; omega model.*

GJCST-A Classification : *C.1.4*



Strictly as per the compliance and regulations of:



A Frame work for Parallel String Matching- A Computational Approach with Omega Model

K Butchi Raju ^α, Chinta Someswara Rao ^σ & Dr. S. Viswanadha Raju ^ρ

Abstract- Now a day's parallel string matching problem is attracted by so many researchers because of the importance in information retrieval systems. While it is very easily stated and many of the simple algorithms perform very well in practice, numerous works have been published on the subject and research is still very active. In this paper we propose a omega parallel computing model for parallel string matching. Experimental results show that, on a multi-processor system, the omega model implementation of the proposed parallel string matching algorithm can reduce string matching time by more than 40%.

Keywords: string matching; parallel string matching; computing model; omega model.

I. INTRODUCTION

String matching has been one of the most extensively studied problems in computer engineering since it performs important tasks in many applications like information retrieval (IRS), web search engines, error correction and several other fields [1-12]. Especially with the introduction of search engines dealing with tremendous amount of textual information presented on the World Wide Web, so this problem deserves special attention and any improvements to speed up the process will benefit these important applications [1-12].

As current free textual databases are growing almost exponentially with the time, the string matching problem becomes impractical to use the fastest sequential algorithms on a conventional sequential computer system [1-12]. To improve the performance of searching on large text collections, some researchers developed special purpose algorithms called parallel algorithms that parallelized the entire database comparison on general purpose parallel computers where each processor performs a number of comparisons independently. In Parallel processing the text string T and pattern P are assumed and that two input words have already been allocated in the processors in such a way that each processor stores a single text symbol, and some processors additionally a single pattern symbol. The input words are stored symbol-by-symbol in consecutive processors numbered according to the snake - like row - major indexing, that

is, the processors in the odd-numbered rows 1, 3, 5, ... are numbered from left to right, and in the even-numbered rows from right to left. (The first symbols of T and P are in processor 1, the next in processor 2, and so on.) This allocation scheme places symbols adjacent in the text or pattern in adjacent processors. The output of the string matching algorithm is that each processor is to be marked as either being a starting position of an occurrence of P in T or not. In this paper we proposed a parallel string matching technique based on butterfly model.

The main contributions of this work are summarized as follows. This work offers a comprehensive study as well as the results of typical parallel string matching algorithms at various aspects and their application on butterfly computing models. This work suggests the most efficient algorithmic butterfly models and demonstrates the performance gain for both synthetic and real data. The rest of this work is organized as, review typical algorithms, algorithmic models and finally conclude the study.

II. RECENT ADVANCEMENTS AND GLOBAL RESEARCH

The first optimal parallel string matching algorithm was proposed by Galil [13]. On SIMD-CRCW model, this algorithm required $n / \log n$ processors, and the time complexity is $O(\log n)$; on SIMD-CREW model, it required $n / \log^2 n$ processors and the time complexity is $O(\log^2 n)$. Vishkin [14] improved this algorithm to ensure it is still optimal when the alphabet size is not fixed. In [15], an algorithm used $O(n \times m)$ processors was presented, and the computation time is $O(\log \log n)$. A parallel KMP string matching algorithm on distributed memory machine was proposed by CHEN [16]. The algorithm is efficient and scalable in the distributed memory environment. Its computation complexity is $O(n / p + m)$, and p is the number of the processors.

SV Raju et.al [17] presents new method for exact string matching algorithm based on layered architecture and two-dimensional array. This has applications such as string databases and computational biology. The main use of this method is to reduce the time spent on comparisons of string matching by distributing the data among processors which achieves a linear speedup and requires layered architecture and additionally $p^* \#$ processors.

Author ^α: Associate Professor, Department of CSE, GRIET, Hyderabad, AP, India.

Author ^σ: Assistant Professor, Dept of CSE, SRKR Engineering College, Bhimavaram, AP, India.

Author ^ρ: Professor & HOD, Department of CSE, JNTU College of Engineering, JNTU Jagithyal, AP, India.

Bi Kun et.al [18] proposed the improved distributed string matching algorithm. And also an improved single string matching algorithm based on a variant Boyer-Moore algorithm is presented. In this they implement algorithm on the above architecture and the experiments prove that it is really practical and efficient on distributed memory machine. Its computation complexity is $O(n/p + m)$, where n is the length of the text, and m is the length of the pattern, and p is the number of the processors. They show that this distributed architecture is suitable for paralleling the multipattern string matching algorithms and approximate string matching algorithms.

Hsi-Chieh Le [19] et.al presents three algorithms for string matching on reconfigurable mesh architectures. Given a text T of length n and a pattern P of length m , the first algorithm finds the exact matching between T and P in $O(1)$ time on a 2-dimensional RMESH of size $(n - m + 1) \times m$. The second algorithm finds the approximate matching between T and P in $O(k)$ time on a 2D RMESH, where k is the maximum edit distance between T and P . The third algorithm allows only the replacement operation in the calculation of the edit distance and finds an approximate matching between T and P in constant-time on a 3D RMESH. By this paper we state that this is simpler model would be sufficient to run the proposed algorithms without increasing the reported time complexities.

S V Raju [20] et.al considers the problem of string matching algorithm based on a two-dimensional mesh. This has applications such as string databases, cellular automata and computational biology. The main use of this method is to reduce the time spent on comparisons in string matching by using mesh connected network which achieves a constant time for mismatch a text string and. This is the first known optimal-time algorithm for pattern matching on meshes. The proposed strategy uses the knowledge from the given algorithm and mesh structure.

Its'hak Dinstein [21] et.al propose a parallel computation approach to two dimensional shape recognition. This approach uses parallel techniques for contour extraction, parallel computation of normalized contour-based feature strings independent of scale and orientation, and parallel string matching algorithms. The implementation on the EREW PRAM architecture is discussed, but it can be adapted to other parallel architectures.

Jin Hwan Park [22] et.al presents efficient dataflow schemes for parallel string matching. In this they consider two sub problems known as the exact matching and the k -mismatches problems are covered. Three parallel algorithms based on multiple input (and output) streams are presented. Time complexities of these parallel algorithms are $O((n/d)+a)$, $0 \leq a \leq m$, where n and m represent lengths of reference and pattern strings ($n \gg m$) and d represents the number

of streams used (the degree of parallelism). They show, they can control the degree of parallelism by using variable number (d) of input (and output) streams. They show their approaches solve the exact matching and the k -mismatches problems with time complexities of $O((n/d) + a)$, where $a = \log m$ for the hierarchical scheme, m for the linear scheme, and 0 for the broadcasting scheme. Required time to process length n reference string is reduced by a factor of d by using d identical computation parts in parallel. With linear systolic array architecture, m PEs are needed for serial design and $d \cdot m$ PEs are needed for parallel design, where m is the pattern size and the d is the controllable degree of the parallelism (i.e. number of streams used).

S V Raju [23] et.al considers the problem of exact string matching algorithm based on a two-dimensional array. This has applications such as string databases, cellular automata and computational biology. The main use of this method is to reduce the time spent on comparisons in string matching by finding common characters in pattern string which achieves a constant time $O(1)$ for pattern string in a text string. This reduces many calls across backend interface.

Chuanpeng Chen [24] et.al propose a high throughput configurable string matching architecture based on Aho-Corasick algorithm. The architecture can be realized by random-access memory (RAM) and basic logic elements instead of designing new dedicated chips. The bit-split technique is used to reduce the RAM size, and the byte-parallel technique is used to boost the throughput of the architecture. By the particular design and comprehensive experiments with 100MHz RAM chips, one piece of the architecture can achieve a throughput of up to 1.6Gbps by 2-byte-parallel input, and we can further boost the throughput by using multiple parallel architectures.

Prasanna[25] et.al propose a multi-core architecture on FPGA to address these challenges. They adopt the popular Aho-Corasick (AC-opt) algorithm for our string matching engine. Utilizing the data access feature in this algorithm, they design a specialized BRAM buffer for the cores to exploit a data reuse existing in such applications. Several design optimizations techniques are utilized to realize a simple design with high clock rate for the string matching engine. An implementation of a 2-core system with one shared BRAM buffer on a Virtex-5 LX155 achieves up to 3.2 GBPS throughput on a 64 MB state transition table stored in DRAM. Performance of systems with more cores is also evaluated for this architecture, and a throughput of over 5.5 Gbps can be obtained for some application scenarios.

S. Muthukrishnan et.al[26] present an algorithm on the CRCW PRAM that checks if there exists a false match in $O(1)$ time using $O(n)$ processors. This algorithm does not require preprocessing the pattern. Therefore, checking for false matches is provably

simpler than string matching since string matching takes $O(\log(\log m))$ time on the CRCW PRAM. In this they use this simple algorithm to convert the Karp–Rabin Monte Carlo type string-matching algorithm into a Las Vegas type algorithm without asymptotic loss in complexity. Finally they present an efficient algorithm for identifying all the false matches and, as a consequence, show that string-matching algorithms take $A \cdot \log \log m /$ time even given the flexibility to output a few false matches.

S V Raju [27] et.al present new approach for parallel string matching. Some known parallel string matching algorithms are considered based on duels by witness who focuses on the strengths and weaknesses of the currently known methods. The new 'divide and conquer' approach has been introduced for parallel string matching, called the W-period, which is used for parallel preprocessing of the pattern and has optimal implementations in a various models of computation. The idea, common for every parallel string matching algorithm is slightly different from sequential ones as Knuth-Morris-Pratt or Boyer-Moore algorithm.

III. COMMUNICATION NETWORK RELATION TO COMPUTER SYSTEM COMPONENTS

A typical distributed system is shown in Figure 1. Each computer has a memory-processing unit and the computers are connected by a communication network. Figure 2 shows the relationships of the software components that run on each of the computers and use the local operating system and network protocol stack for functioning.

The distributed software is also termed as middleware. A distributed execution is the execution of processes across the distributed system to collaboratively achieve a common goal. An execution is also sometimes termed a computation or a run. The distributed system uses a layered architecture to break down the complexity of system design. The middleware is the distributed software that drives the distributed system, while providing transparency of heterogeneity at the platform level [28]. Figure 2 schematically shows the interaction of this software with these system components at each processor.

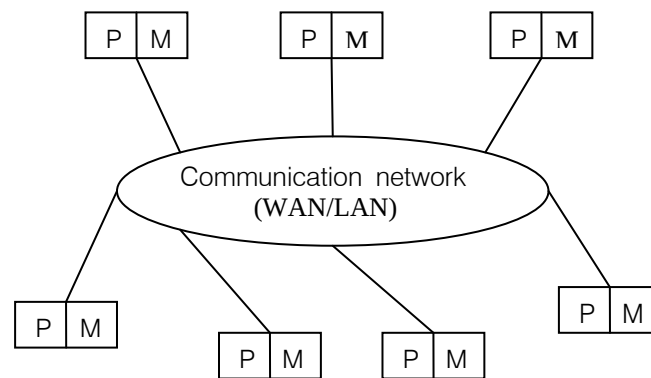


Figure 1 : A distributed systems connects processors by a communication network.

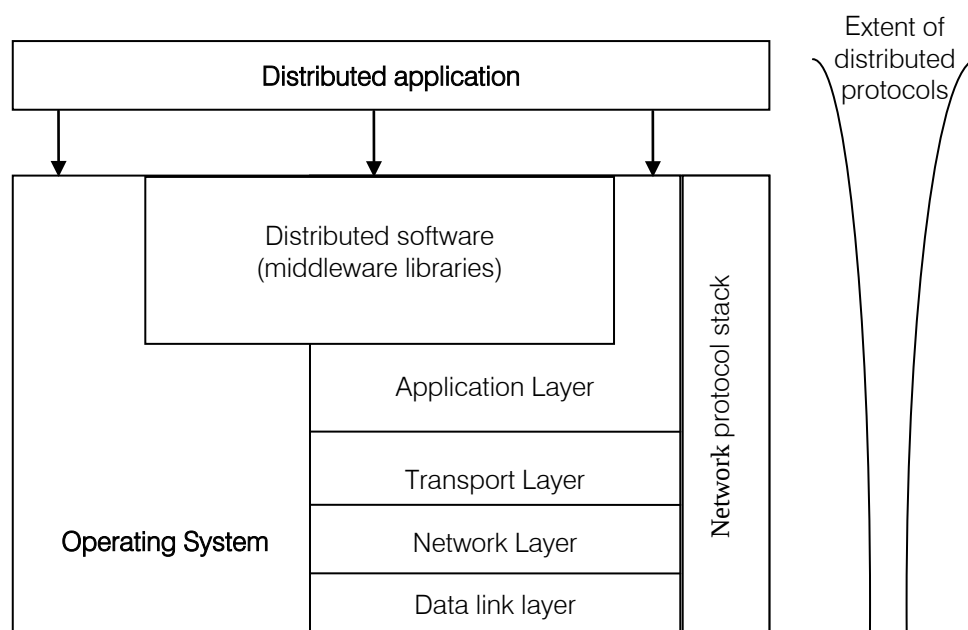


Figure 2 : Interaction of the software components at each processor.

IV. TEXT PARTITIONING

The exact string-matching problem can achieve data parallelism with data partitioning technique. We decompose the text into r subtexts, where each subtext contains $(T/p)+m-1$ successive characters of the complete text. There is an overlap of $m-1$ string characters between successive subtexts, i.e, a redundancy of $r(m-1)$ characters. Alternatively it could be assumed that the database of an information retrieval system contains r independent documents. Therefore, in both the cases all the above partitions yield a number of independent tasks each comprising some data (i.e. a string and a large subtext) and a sequential string matching procedure that operates on that data. Further, each task completes its string matching operation on its local data and returns the number of occurrences[29-31]. Finally, we can observe that there are no communication requirements among the tasks but only global (or collective) communication is required.

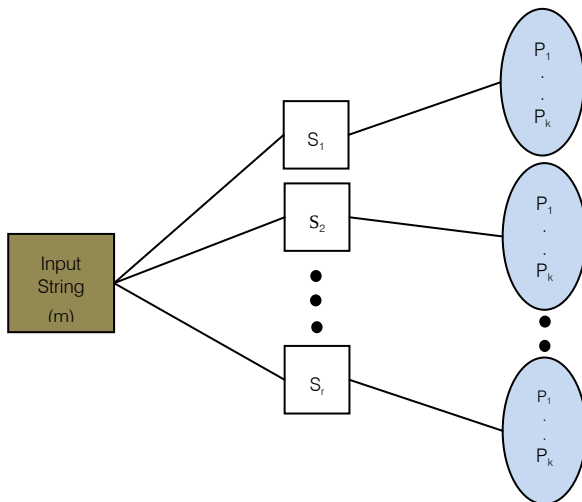


Figure 3 : Framework for pool of Processors

The main issue to be addressed is how the several tasks (or r subtexts) can be mapped or distributed to multiple processors for concurrent execution. In [29-31] different ways of distributing the database across a multi computer network were discussed. Let p be the number of processors in network and r be the number of subtext in the whole collection then the text partition is defined as, if $r=p$ then each subtext contains $T/p+m-1$ characters. This is called static allocation of subtext as shown in Fig 3. In the next section we present the parallel algorithm that is based on static allocation of subtext using MPI library. A significant contribution of this paper is a demonstration of the maximum size buffer with $2k$ processors for implementation of string matching and capable of accepting a character from r subtexts where $k=8$ bits. This architecture enables a buffered string matching system implementing a KMP like pre computation algorithm. In the above mapping $\{a_1, a_2...a_r\}$ is the

input string where r represents subtext and $\{p_1, p_2...p_k\}$ are the number of processors for the given input string ($k=8$). In the above mapping the given input string will be allocated to each processor as shown in Fig. 3.

V. METHODOLOGY

Multistage interconnection networks (MINs) consist of more than one stages of small interconnection elements called switching elements and links interconnecting them. Multistage interconnection networks (MINs) are used in multiprocessing systems to provide cost-effective, high-bandwidth communication between processors and/or memory modules. A MIN normally connects N inputs to N outputs and is referred as an $N \times N$ MIN. The parameter N is called the size of the network[32-33]. The popularity of MINs stems from both the operational features they deliver –e.g. their ability to route multiple communication tasks concurrently- and the appealing cost/performance ratio they achieve. MINs with the Banyan [34] property e.g. Omega Networks [35], Delta Networks [36], and Generalized Cube Networks [37] are more widely adopted, since non-Banyan MINs have -generally- higher cost and complexity. Both in the context of parallel and distributed system, the performance of the communication network interconnecting the system elements (nodes, processors, memory modules etc) is recognized as a critical factor for overall system performance. Consequently, the need for communication infrastructure performance prediction and evaluation has arisen, and numerous research efforts have targeted this area, employing either analytical models (mainly based on Markov models and Petri-nets) or simulation techniques.

There are several different multistage interconnection networks proposed and studied in the literature. Figure1 illustrates a structure of multistage interconnection network, which are representatives of a general class of networks. This Figure 4 shows the connection between p inputs and b outputs, and connection between these is via n number of stages.

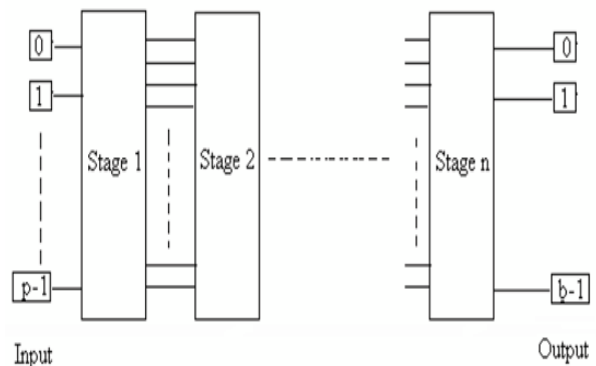


Figure 4 : A Multistage Interconnection Network(MIN)

A multistage interconnection network is actually a compromise between crossbar and shared bus networks multistage interconnection networks are:

- Attempt to reduce cost
- Attempt to decrease the path length

In a multistage interconnection network, as in a crossbar, switching elements are distinct from processors. Instead messages pass through a series of switch stages. The network can be constructed from unidirectional or bi-directional switches and links. In a unidirectional MIN, all messages must traverse the same number of wires, and so the cost of sending a message is independent of processor location. In effect, all processors are equidistant. In a bidirectional MIN, the number of wires traversed depends to some extent on processor location, although to a lesser extent than a mesh or hypercube.

VI. OMEGA NETWORK

Omega network connecting P processors to P memory banks as shown in Fig 5. In general, it consists of $p = (q + 1)2^q$ processors, organized as $q + 1$ ranks of 2^q processors each (Figure 1). Optionally we shall identify the rightmost and the leftmost ranks, so there is no rank q , and the processors on ranks 0 and $q - 1$ are connected directly. Let us denote the processor i on the rank r by $P_{ir,r}$, $0 \leq i < 2^q$, $0 \leq r \leq q$. Then processor $P_{ir,r+1}$ is connected to the two processors $P_{ir,r}$ and $P_{ir,r+1}$ and processor $P_{ir,r+1}$ is connected to the two processors $P_{ir,r}$ and $P_{ir,r+1}$. Recall that $i' = i_{q-1} \dots i_1 i_0$. These four connections form a "butterfly" pattern, from which the name of the network is derived. The hypercube is actually the butterfly with the rows collapsed. The communication link in the hypercube between processors P_{i_1} and P_{i_2} is identified with the communication links in the butterfly between $P_{i_1,r+1}$ and $P_{i_2,r+1}$, and between $P_{i_1,r+1}$ and $P_{i_2,r}$.

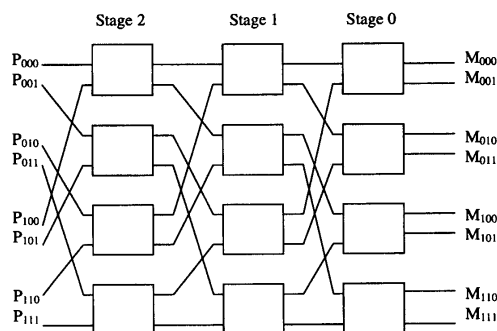


Figure 5 : Omega Network

Terminology : Terminal reliability is defined as the probability of successful communication between an input output pair. In this section, terminal reliability of Omega, has evaluated. The Omega is a unique-path MIN that has N input switches and N output switches and n stages, where $n = \log_2 N$. An 8×8 Omega has

three stages, 12 SEs and 32 links. And reliability is shown in Fig 6. Let r be the probability of a switch being operational. As Omega is a unique-path MIN, the failure of any switch will cause system failure, so from the reliability point of view, there are $\log_2 N$ SEs in series for each terminal path. Hence, the terminal reliability of an $N \times N$ Omega is $R_t(\Omega) = (r)^{\log_2 N}$. As there is only a single path between a particular input S_i , $i = 1, 2, 3, 4$, and an output in an 8×8 Omega so the terminal reliability is $R_t(\Omega) = (r)^3$.

Switching Reliability	Terminal Reliability of Omega
0.99	0.970299
0.98	0.941192
0.96	0.884736
0.95	0.857375
0.94	0.830584
0.92	0.778688
0.9	0.729

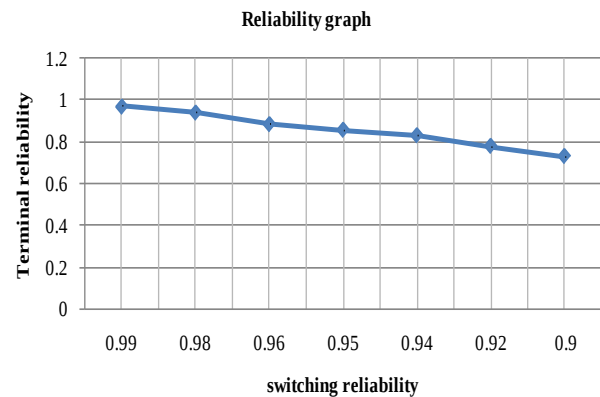


Figure 6 : Switching Reliability

VII. PROPOSED SYSTEM STRUCTURE

In this we propose a system for parallel processing with omega model. Its shared-memory, expandable MIMD parallel computer.

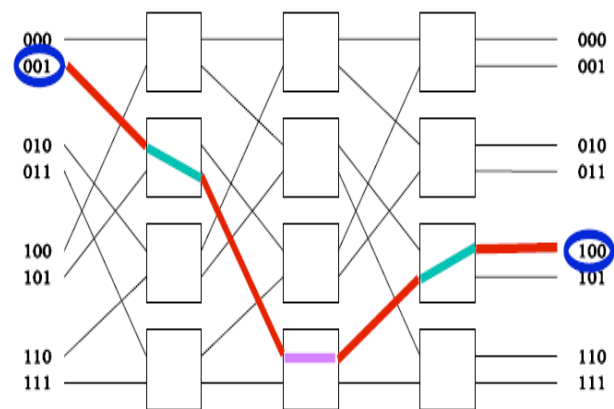


Figure 7 : Omega Network with 8 i/o

The computer got its name from the omega switch which it uses for interprocessor communication.

The switch supports a processor-to-processor bandwidth of 32 Mbits/second. Figure 7 illustrates 8-input 8-output omega switch.

VIII. PROGRAMMING THE OMEGA PARALLEL PROCESSOR

The omega Parallel Processor is programmed exclusively in high-level programming languages. Searching, IRS, Editing, compiling and linking, downloading, running and debugging of programs are done from a UNIX front-end. A window manager enables rapid switching between the front-end and the Butterfly system environments. Two distinct approaches to programming the omega have seen widespread use: message passing and shared memory. When using the message passing paradigm the programmer decomposes the application into a moderately sized collection of loosely coupled processes which from time to time exchange control signals or data. This approach is similar to programming a multiprocessor application for a uniprocessor. In the shared memory approach, a

task is usually some small procedure to be applied to a subset of the shared memory. A task, therefore, can be represented simply as an index, or a range of indices, into the shared memory and an operation to be performed on that memory. This style is particularly effective for applications containing a few frequently repeated tasks. Memory and processor management are used to keep all memories and processors equally busy.

IX. PARALLEL STRING MATCHING ALGORITHM ON OMEGA MODEL

In this model data on processors have been organized such that they represent the m sets of length of $n-m+1$ of the text string with $m^* n-m+1$ matrix plus, the first processor of each row segment holding the first element of each set also carries an element of pattern. The process is similar as per above for the remaining $m-1$ rows. First show how to find the occurrences of pattern P in text string T on omega model with $m^*(n-m+1)$ in constant time $O(1)$.

Algorithm for string matching (pattern P , text T)	
Input: m elements of pattern initially distributed to the m processors on the first column, one processor and $L_{i,j}$ distributed to the $m^*(n-m+1)$ processors on the row and column. Where $1 \leq i \leq m$ and $m \leq j \leq m^*(n-m+1)$	
Starting processors passes elements on row buses: each first processor $p_{i,1}$, where $1 \leq i \leq m$, broadcasts the element $c_{i,1}$ to every processor in the i th row using only row buses. After this multiple row broadcast communication operations each processor $p_{i,j}$ saves received element as $r_{i,j}$.	
Compare and Set results: Each processor $p_{i,j}$ compares $r_{i,j}$ with $X_{i,j}$, if $r_{i,j} = X_{i,j}$ if sets result=1 otherwise, result=0	
Sum up 1's: perform a one-dimensional binary prefix sum operation on each column simultaneously for the value of result. Each processor $p_{i,j}$ where $1 \leq (i,j) \leq m$ stores the binary prefix sums to $b_{i,j}$.	
Distance: If $P_{m,m} = 0$ then the string matching(i.e. exact string matching) otherwise approximate string matching with k mismatches.	

Lemma-1 : Each step in the above algorithm runs in constant time. Thus we have the following theorem.

Theorem 1.1 : There is a constant time string matching algorithm on a omega model that finds the occurrences of pattern in text using $m^*(n-m+1)$ processors

Example : Text string $T(n) = \text{GOKARAJU}$ and Pattern string $P(m) = \text{RAJU}$ $L1 = \{\text{GOKAR}\}$, $L2 = \{\text{OKARA}\}$ $L3 = \{\text{KARAJ}\}$ $L4 = \{\text{ARAJU}\}$

Step-1

R	G	O	K	A	R
A	O	K	A	R	A
J	K	A	R	A	J
U	A	R	A	J	U

Step-2

<R,G>	<R,O>	<R,K>	<R,A>	<R,R>
<A,O>	<A,K>	<A,A>	<A,R>	<A,A>
<J,K>	<J,A>	<J,R>	<J,A>	<J,J>
<U,A>	<U,R>	<U,A>	<U,J>	<U,U>

Step-3

0	0	0	0	1
0	0	1	0	1
0	0	0	0	1
0	0	0	0	1

Step-4

0	0	0	0	1
0	0	1	0	1
0	0	0	0	1
0	0	0	0	1
4	4	3	4	0

As per the given example, after step 4 in the matrix $M_{m+1,j}$ values useful for deciding the matching is exact string matching or approximate string matching with the k mismatches.

Lemma-2 : So that the string matching is completely scalability and obtain the following theorem.

Theorem 1.2 : The given two strings size of text n and size of pattern m . find the occurrences of pattern in text.

There is completely scalable on Butterfly model. The algorithm runs in $O(m*(n-m+1)/P)$ time, where P is the number of processors and $1 \leq p \leq m*(n-m+1)$.

X. PERFORMANCE EVALUATION

In order to evaluate the overall performance of a multi-priority ($\Lambda \& M$) MIN consisting of (2×2) SEs, we use the following metrics. Let T be a relatively large time period divided into u discrete time intervals $(\tau_1, \tau_2, \dots, \tau_u)$.

Average throughput the average number of packets accepted by all destinations per network cycle. Formally, Th_{avg} (or *bandwidth*) is defined as

$$Th_{avg} = \lim_{u \rightarrow \infty} \frac{\sum_{k=1}^u n(k)}{u} \quad (1)$$

where $n(k)$ denotes the number of packets that reach their destinations during the k th time interval.

Normalized throughput is the ratio of the average throughput Th_{avg} to number of network outputs N . Formally, Th can be expressed by back ground and reflects how effectively network capacity is used.

$$Th = \frac{Th_{avg}}{N} \quad (2)$$

Relative normalized throughput $RTh(i)$ of i -class priority traffic, where $i=1..p$ is the *normalized throughput* $Th(i)$ of i -class priority packets divided by the corresponding-class *offered load* $\lambda(i)$ of such packets.

$$RTh(i) = \frac{Th(i)}{\lambda(i)} \quad (3)$$

Average packet delay $Davg(i)$ of i -class priority traffic, where $i=1..p$ is the average time a corresponding-class priority packet spends to pass through the network. Formally, $Davg(i)$ is expressed by

$$D_{avg}(i) = \lim_{u \rightarrow \infty} \frac{\sum_{k=1}^{n(u)} t_d(k)}{n(u)} \quad (4)$$

where $n(u)$ denotes the total number of the corresponding- class priority packets accepted within u time intervals and $td(k)$ represents the total delay for the k th such packet. We consider $td(k) = tw(k) + ttr(k)$ where $tw(k)$ denotes the total queuing delay for k th packet waiting at each stage for the availability of a

corresponding-class empty buffer at the next stage queue of the network.

The second term $ttr(k)$ denotes the total transmission delay for k th such packet at each stage of the network, that is just $n*nc$, where $n=\log_2 N$ is the number of intermediate stages and nc is the network cycle.

XI. CONCLUSIONS

In this paper we concentrate on parallel algorithms for string matching on computing models, especially in omega model. In this paper simulate the parallel algorithms for the implementation of high speed string matching; this uses fine-grained parallelism and performs matching of a search string by splitting the string into a set of substrings and then matching all of the substrings simultaneously. We also see that this implementation can be optimized in terms of resource utilization.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Chinta Someswararao, K Butchiraju, S ViswanadhaRaju, "Recent Advancement is Parallel Algorithms for String matching on computing models - A survey and experimental results", LNCS, Springer, pp.270-278, ISBN: 978-3-642-29279-8 ,2011.
2. Chinta Someswararao, K Butchiraju, S ViswanadhaRaju, "PDM data classification from STEP- an object oriented String matching approach", IEEE conference on Application of Information and Communication Technologies, pp.1-9, ISBN: 978-1-61284-831-0, 2011.
3. Chinta Someswararao, K Butchiraju, S ViswanadhaRaju, "Recent Advancement is Parallel Algorithms for String matching - A survey and experimental results", IJAC, Vol 4 issue 4, pp-91-97, 2012.
4. Simon Y. and Inayatullah M., "Improving Approximate Matching Capabilities for Meta Map Transfer Applications," Proceedings of Symposium on Principles and Practice of Programming in Java, pp.143-147, 2004.
5. Chinta Someswararao, K Butchiraju, S ViswanadhaRaju, "Parallel Algorithms for String Matching Problem based on Butterfly Model", pp.41-56, IJCST, Vol. 3, Issue 3, July – Sept, ISSN 2229-4333, 2012.
6. Chinta Someswararao, K Butchiraju, S ViswanadhaRaju, "Recent Advancement is String matching algorithms- A survey and experimental results", IJCIS, Vol 6 No 3, pp.56-61, 2013.
7. Chinta Someswararao, " Parallel String Matching Problems with Computing Models - An Analysis of the Most Recent Studies", International Journal of Computer Applications , Vol.76(15), pp.7-25,

- Published by Foundation of Computer Science, New York, USA, 2013.
8. Chinta Someswararao, "Parallel String Matching with Multi Core Processors-A Comparative Study for Gene Sequences", *Global Journal of Computer Science and Technology*, Vol-13, Issue-1, pp.27-41, 2013.
 9. K, Grabowski S, "Average-Optimal String Matching", *Journal of Discrete Algorithms*, pp- 579-594,2009.
 10. Luis Russo L, Navarro G, Oliveira A, Morales P, "Approximate String Matching with Compressed Indexes Algorithm", pp- 1105-1136,2009.
 11. Ilie L, Navarro G, Tinta L, "The Longest Common Extension Problem, Revisited and Applications to Approximate String Searching", *Journal of Discrete Algorithms*, pp-418-428, 2010.
 12. Fredriksson K, Grabowski S, "Average-Optimal String Matching, *Journal of Discrete Algorithms*", pp- 579-594,2009.
 13. Z. Galil, "Optimal parallel algorithms for string matching," in *Proc. 16th Annu. ACM symposium on Theory of computing*, pp. 240-248, 1984.
 14. U. Vishkin, "Optimal parallel matching in strings," *Information and control*, vol. 67, pp. 91-113, 1985.
 15. Y. Takefuji, T. Tanaka, and K. C. Lee, "A parallel string search algorithm", *IEEE Trans. Systems, Man and Cybernetics*, vol. 22, pp. 332-336, March-April 1992.
 16. CHEN Guo-liang, LIN-Jie, and GU Nai-jie, "Design and analysis of string matching algorithm on distributed memory machine," *Journal of Software*, vol. 11, pp. 771-778, 2000.
 17. Viswanadha Raju, S.; Vinaya Babu, A.; Mrudula, M.; "Backend Engine for Parallel String Matching Using Boolean Matrix", *IEEE on PAR ELEC*, pp-281-283,2006.
 18. Bi Kun, Gu Nai-jie, Tu Kun, Liu Xiao-hu, and Liu Gang A Practical Distributed String Matching Algorithm Architecture and Implementation World Academy of Science, Engineering and Technology , 2005.
 19. Hsi-Chieh Leet,Fikret Ercalt, "RMESH Algorithms For Parallel String Matching" ,IEEE,1997.
 20. S. Viswanadha Raju and A. Vinayababu, "Optimal Parallel algorithm for String Matching on Mesh Network Structure", 2006.
 21. Its'hak Dinstein, ad M. Landau, "Using Parallel String Matching Algorithms for Contour Based 2-D Shape Recognition", IEEE,1990.
 22. Jin Hwan Park and K. M. George, "Parallel String Matching Algorithms Based on Dataflow", *IEEE on System Sciences*, 1999.
 23. S.Viswanadha Raju S R Mantena A.Vinaya Babu G V S Raju, "Efficient Parallel Pattern Matching using Partition Method",2006.
 24. Chuanpeng Chen, Zhongping Qin, "A Bit-split Byte-parallel String Matching Architecture",IEEE,2009.
 25. Qingbo Wang, Viktor K. Prasanna, "Multi-Core Architecture on FPGA for Large Dictionary String Matching", *IEEE on Field Programmable Custom Computing Machines*,2009.
 26. S. Muthukrishnan "Detecting False Matches in String-Matching Algorithms", *Algorithmica* ,Springer-Verlag New York Inc.1997.
 27. S.Viswanadha Raju , A.Vinaya Babu, G.V.S.Raju, and K.R. Madhavi , "W-Period Technique for Parallel String Matching",2007.
 28. Ajay D. Kshemkalyani and Mukesh Singhal, "Distributed Computing: Principles, Algorithms, and Systems",Cambridge.
 29. S.Viswanadha Raju and A.Vinayababu, 2004, "Performance in the design of Parallel Programming", *Proc ObComAPC-2004*, Allied Publications, pp.380 to 392.
 30. S.Viswanadha Raju, A.Vinayababu, S.P.Yanaiah and GVS Raju, 2006 "Parallel Approach for K String Matching", *Proc NCIMDiL-2006*, Indian Institute Of Technology, Kharagpur , 5-10.
 31. J. Garofalakis, and E. Stergiou "An analytical performance model for multistage interconnection networks with blocking", *Procs. of CNSR 2008*,May 2008
 32. Josep Torrellas, Zheng Zhang. The Performance of the Cedar Multistage Switching Network. *IEEE Transactions on Parallel and Distributed Systems*, 8(4), pp. 321-336, 1997.
 33. Bhogavilli S. K., Abu-Amara H., "Design and Analysis of High-performance Multistage Interconnection Networks", *IEEE Transactions on Computers*, vol. 46, no. 1, January 1997, pp. 110 - 117.
 34. G. F. Goke, G.J. Lipovski. "Banyan Networks for Partitioning Multiprocessor Systems" *Procs. of 1st Annual Symposium on Computer Architecture*, pp. 21-28, 1973.
 35. D. A. Lawrie. "Access and alignment of data in an array processor", *IEEE Transactions on Computers*, C-24(12):1145-1155,Dec. 1975.
 36. J.H. Patel. "Processor-memory interconnections for mutliprocessors", *Procs. of 6th Annual Symposium on Computer Architecture*. New York, pp. 168-177, 1979.
 37. G. B. Adams and H. J. Siegel, "The extra stage cube: A fault-tolerant interconnection network for supersystems", *IEEE Trans. on Computers*, 31(4)5, pp. 443-454, May 1982.



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 2 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

A Performance Comparison Between Enlightenment and Emulation in Microsoft Hyper-V

By Hasan Fayyad-Kazan, Luc Perneel & Martin Timmerman

Vrije Universiteit Brussel, Belgium

Abstract- Microsoft (MS) Hyper-V is a native hypervisor that enables platform virtualization on x86-64 systems. It is a micro-kernelized hypervisor where a host operating system provides the drivers for the hardware. This approach leverages MS Hyper-V to support enlightenments (the Microsoft name for Paravirtualization) in addition to the hardware emulation virtualization technique.

This paper provides a quantitative performance comparison, using different tests and scenarios, between enlightened and emulated Virtual Machines (VMs) hosted by MS Hyper-V server 2012. The experimental results show that MS enlightenments improve performance by a factor of more than two.

Keywords: *virtualization, hyper-v, enlightenments, emulation.*

GJCST-A Classification : *C.1.4*



A PERFORMANCE COMPARISON BETWEEN ENLIGHTENMENT AND EMULATION IN MICROSOFT HYPER-V

Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

A Performance Comparison Between Enlightenment and Emulation in Microsoft Hyper-V

Hasan Fayyad-Kazan ^α, Luc Perneel ^σ & Martin Timmerman ^ρ

Abstract- Microsoft (MS) Hyper-V is a native hypervisor that enables platform virtualization on x86-64 systems. It is a micro-kernelized hypervisor where a host operating system provides the drivers for the hardware. This approach leverages MS Hyper-V to support enlightenments (the Microsoft name for Paravirtualization) in addition to the hardware emulation virtualization technique.

This paper provides a quantitative performance comparison, using different tests and scenarios, between enlightened and emulated Virtual Machines (VMs) hosted by MS Hyper-V server 2012. The experimental results show that MS enlightenments improve performance by a factor of more than two.

Keywords: virtualization, hyper-v, enlightenments, emulation.

I. INTRODUCTION

Virtualization has become a popular way to make more efficient use of server resources within both private data centers and public cloud platforms. It refers to the creation of a Virtual Machine (VM) which acts as a real computer with an operating system (OS) [1]. It also allows sharing the underlying physical machine resources with different VMs.

The software layer providing the virtualization is called a Virtual Machine Monitor (VMM) or hypervisor [1]. It can be either Type 1 (or native, bare metal) running directly on the host's hardware to control the hardware and to manage guest operating systems, or Type 2 (or hosted) running within a conventional operating-system environment.

Since it has direct access to the hardware resources rather than going through an operating system, a native hypervisor is more efficient than a hosted architecture and delivers greater scalability, robustness and performance [2].

Microsoft Hyper-V implements Type 1 hypervisor virtualization [3]. In this approach, a hypervisor runs directly on the hardware of the host system and is

responsible for sharing the physical hardware resources with multiple virtual machines [4]. In basic terms, the primary purpose of the hypervisor is to manage the physical CPU(s) and memory allocation between the various virtual machines running on the host system.

There are several ways to implement virtualization. Two leading approaches are Full virtualization (FV)/Hardware emulation and Para-virtualization (PV) [5]. Enlightenment is the Microsoft name for Para-virtualization.

This paper provides a quantitative performance comparison between hardware emulation and Enlightenments (Para-Virtualization) techniques hosted by MS Hyper-V server 2012.

It is organized as follows: Section 2 describes MS Hyper-V architecture and Enlightenment approach; Section 3 shows the experimental setup used for our evaluation; Section 4 explains the test metrics, scenarios and results obtained; and section 5 gives a final conclusion.

II. MICROSOFT HYPER-V

Microsoft Hyper-V is a hypervisor-based virtualization technology for x64 versions of Windows Server [6]. It exists in two variants: as a stand-alone product called Hyper-V Server and as an installable role/component in Windows Server [7].

There is no difference between MS Hyper-V in each of these two variants. The hypervisor is the same regardless of the installed edition [7].

MS Hyper-V requires a processor with hardware-assisted virtualization functionality, enabling a much more compact virtualization codebase and associated performance improvements [3].

The Hyper-V architecture is based on microkernelized hypervisors (figure 1). This is an approach where a host operating system, referred to as the parent partition, provides management features and the drivers for the hardware [8].

Author α: PhD Candidate, Department of Electronics and Informatics, Vrije Universiteit Brussel, Pleinlaan 2- 1050 Brussels, Belgium.
e-mail: hafayyad@vub.ac.be

Author σ: PhD Candidate, Department of Electronics and Informatics, Vrije Universiteit Brussel, Pleinlaan 2- 1050 Brussels, Belgium.
e-mail: luc.perneel@vub.ac.be

Author ρ: Professor, Department of Electronics and Informatics, Vrije Universiteit Brussel, Pleinlaan 2- 1050 Brussels, Belgium.
e-mail: martin.timmerman@vub.ac.be

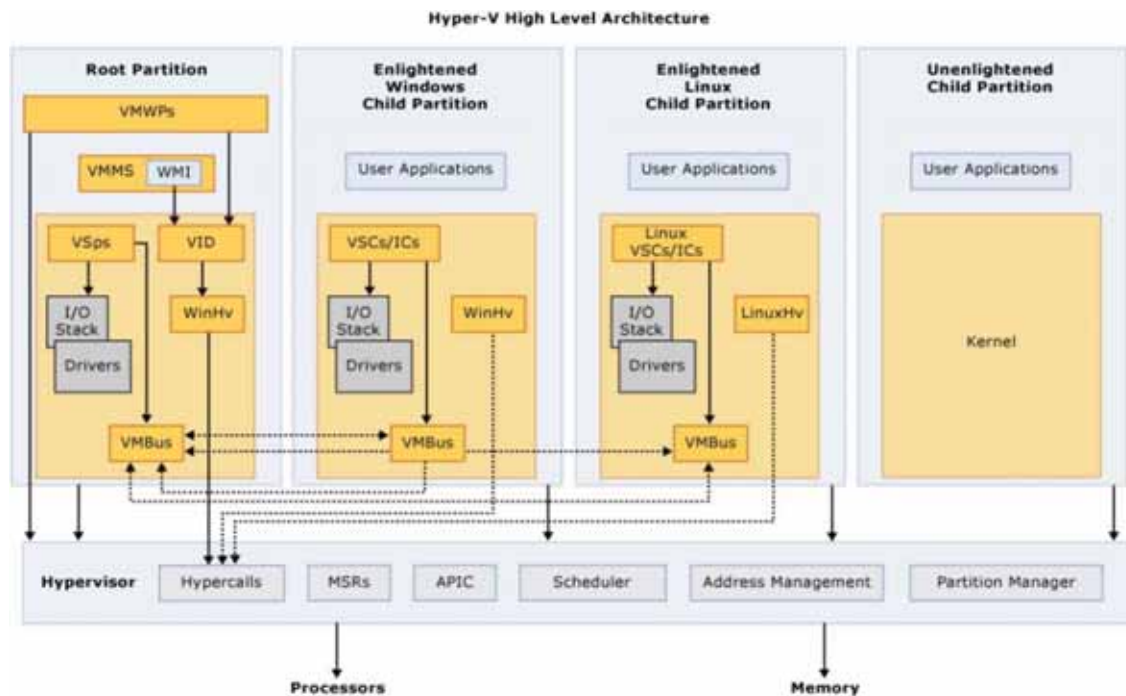


Figure 1 : Hyper-V Architecture [6]

With this approach, the only layer between a guest operating system and the hardware is a streamlined hypervisor with simple partitioning functionality. The hypervisor has no third-party device drivers [9]. The drivers required for hardware sharing reside in the host operating system, which provides access to the rich set of drivers already built for Windows [9].

MS Hyper-V implements isolation of virtual machines in terms of a partition (operating system and applications). A hypervisor instance has to have at least one parent partition, running a supported version of Windows Server [6]. The virtualization stack runs in the parent partition and has direct access to the hardware devices. The parent partition then creates the child partitions which host the guest OSs [6].

Child partitions do not have direct access to hardware resources. Hyper-V can host two categories of operating systems in the child partitions: Enlightened (Hyper-V Aware) and un-enlightened (Hyper-V Unaware) operating systems [10]. Enlightened partition has a virtual view of the resources, in terms of virtual devices. Any request to the virtual devices is redirected via the VMBus (figure 1) – a logical channel which enables inter-partition communication - to the devices in the parent partition managing the requests. Parent partitions run a Virtualization Service Provider (VSP), which connects to the VMBus and handles device access requests from child partitions [6]. Enlightened child partition virtual devices internally run a Virtualization Service Client (VSC) (figure1), which redirect the request to VSPs in the parent partition via the VMBus [6]. The

VSCs are the drivers of the virtual machine, which together with other integration components are referred to as Enlightenments that provide advanced features and performance for a virtual machine. In contrast, the unenlightened child partition does not have the integration components and the VSCs; everything is emulated.

III. EXPERIMENTAL SETUP

Microsoft Hyper-V Server 2012 is tested here. It is a dedicated stand-alone product that contains the hypervisor, Windows Server driver model, virtualization capabilities, and supporting components such as failover clustering, but does not contain the robust set of features and roles found in the Windows Server operating system [11].

As MS Hyper-V supports enlightened and emulated VMs, both VMs are created, running Linux PREEMPTRT v3.8.4-rt2 [12]. Being open source and configurable for usage in enlightened VM are the main reasons for selecting it as the guest OS. This also permits us to compare with another ongoing study of XEN.

Both Linux versions (for enlightened and emulated VM) are built using the buildroot [13] tool to make sure that the enlightenment drivers are added to the enlightened VM.

The tests are done in each VM separately. Under Test VM (UTVM) is the name used for the tested VM, which can be either enlightened or emulated. Each VM has one virtual CPU (vCPU).

One physical CPU is allocated for each VM, using the “virtual machine reserve” and “virtual machine limit” attributes in the VM settings using Hyper-V Manager.

The hardware platform used for conducting the tests has the following characteristics: Intel® Desktop Board DH77KC, Intel® Xeon® Processor E3-1220v2 with 4 cores each running at a frequency of 3.1 GHz, and no hyper-threading support. The cache memory size is as follows: each core has 32 KB of L1 data cache, 32KB of L1 instruction cache and 256 KB of L2 cache. L3 cache is 8MB accessible to all cores. The system memory is 8 GB.

IV. TESTING PROCEDURES AND RESULTS

a) Measuring Process

The Time Stamp Counter (TSC) is used for obtaining (tracing) the measurement values. It is a 64-bit register present on all x86 processors since the Pentium. The instruction RDTSC is used to return the TSC value. This counting register provides an excellent high-resolution, low-overhead way of getting CPU timing information and runs at a constant rate.

b) Testing Metrics

Below is an explanation of the evaluation tests. Note that the tests are initially done on a non-virtualized machine (Bare-Machine) as a reference, using the same OS of the UTM.

i. Clock tick processing duration

The kernel clock tick processing duration is examined here. The results of this test are extremely important as the clock interrupt - being on a high level interrupt on the used hardware platform - will bias all other performed measurements. Using a tickless kernel will not prevent this from happening as it will only lower the number of occurrences. The kernel is not using the tickless timer option.

Here is a description of how this test is performed: a real-time thread with the highest priority is created. This thread does a finite loop of the following tasks: starting the measurement by reading the time using RDTSC instruction, executing a “busy loop” that does some calculations and stopping the measurement by reading the time again using the same instruction. Having the time before and after the “busy loop” provides the time needed to finish its job. In case we run this test on the bare-machine, this “busy loop” will be delayed only by interrupt handlers. As we remove all other interrupt sources, only the clock tick timer interrupt can delay the “busy loop”. When the “busy loop” is interrupted, its execution time increases.

Running the same test in a VM also shows when it is scheduled away by the VMM, which in turn impacts latency.

Figure 2 presents the results of this test on the baremachine, followed by an explanation. The X-axis

indicates the time when a measurement sample is taken with reference to the start of the test. The Yaxis indicates the duration of the measured event; in this case the total duration of the “busy loop”.

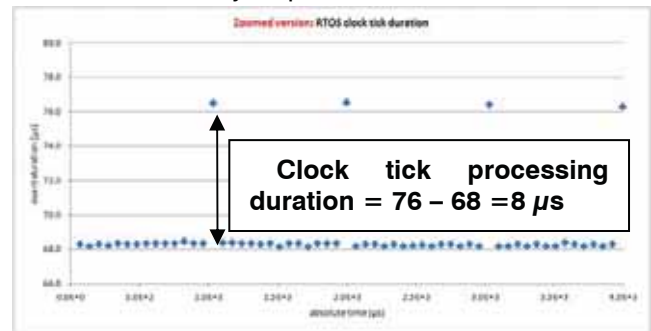


Figure 2 : Clock tick processing duration of the baremachine-zoomed

The lower values (68 μs) of figure 2 present the “busy loop” execution durations if no clock tick happens. In case of clock tick interruption, its execution is delayed until the clock interrupt is handled, which is 76 μs (top values). The difference between the two values is the delay spent handling the tick (executing the handler), which is 8 μs.

Note that the kernel clock is configured to run at 1000 Hz, which corresponds to a tick each 1 ms. This is obvious in figure 2, which is a zoomed version of figure 3 below.

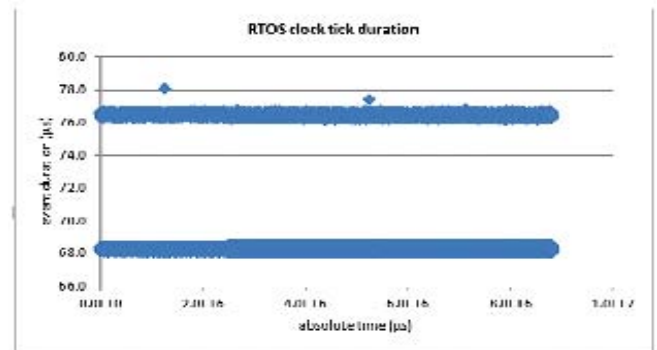


Figure 3 : Clock tick processing duration of the baremachine

Figure 3 represents the test results of 128000 captured samples, in a time frame of 9 seconds. Due to scaling reasons, the samples form a line. As shown in figure 3, the “busy loop” execution time is 78 μs at some periods. Therefore, a clock tick delays any task by 8 to 10 μs.

This test is very useful as it detects all the delays that may occur in a system during runtime. Therefore, we execute this test for long duration (more than one hour) to capture 50 million samples. The results in the tables of section D are comparing the maximum results obtained from the 50 million samples.

ii. Thread switch latency between threads of same priority

This test measures the time needed to switch between threads having the same priority. Although real-time threads should be on different priority levels to be capable of applying rate monotonic scheduling theory [16], this test is executed with threads on the same priority level in order to easily measure thread switch latency without interference of something else.

For this test, threads must voluntarily yield the processor for other threads, so the SCHED_FIFO scheduling policy is used. If we didn't use the FIFO policy, a round-robin clock event could occur between the yield and the trace, and then the thread activation would not be seen in the test trace. The test looks for worst-case behavior and therefore it is done with an increasing number of threads, starting with 2 and going up to 1000. As we increase the number of active threads, the caching effect becomes visible as the thread context will no longer be able to reside in the cache.

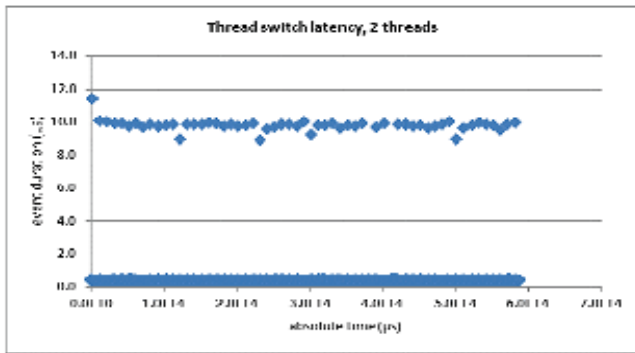


Figure 4: Thread switch latency between 2 threads on the Bare-machine

Figure 4 shows that the minimum switch latency between 2 threads is around $0.43 \mu\text{s}$; the maximum latency is $11.2 \mu\text{s}$ which is dependent on the clock tick processing duration.

Table 1 below shows the results of performing this test on the bare-machine using 2 and 1000 threads.

Table 1: Comparing the "Thread switch latency" results for the bare-machine

Test	Maximum
Switch latency between 2 threads	$11.2 \mu\text{s}$
Switch latency between 1000 threads	$11.45 \mu\text{s}$

Both tests, "Clock tick processing duration" and "Threads switch latency" are done on the enlightened and emulated VMs, using the several scenarios described in section D.

c) Processor Affinity in MS Hyper-V

Most virtualization solutions like Xen and VMware support the affinity concept where a vCPU of a VM can be tied to a given physical processor. Benjamin

Armstrong, Hyper-V Program Manager, explains in the blog "Processor Affinity and why you do not need it on Hyper-V" [14] that there is no need for this concept in Hyper-V. Instead, one can reserve a physical CPU (pCPU) for the VM to guarantee that it always has a whole processor.

Moreover, if a VM has one vCPU and the host has more than one core, this VM can be mapped to any of the available cores in a round-robin way between all the cores [15]. The parent partition is the only VM parked on core 0.

d) Testing scenarios

Below is a description of the scenarios used for the evaluation. In all the scenarios drawings, the parent partition (VM) is not shown because it is idle.

i. Scenario 1: One-to-All

As shown in figure 5, this scenario has only one VM, the UTVM with one vCPU. This vCPU can run on any core (physical CPU) during runtime. The aim of this scenario is to detect the pure hypervisor overhead (as there is no contention).

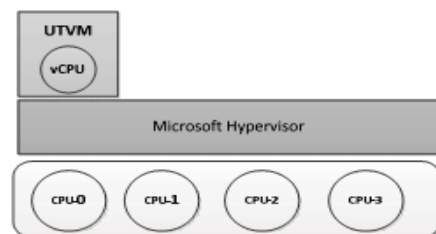


Figure 5: One-to-All scenario

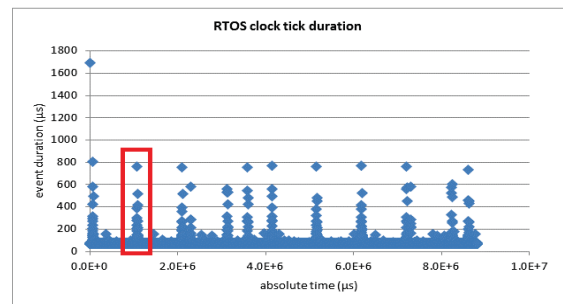


Figure 6: Clock tick duration test for Enlightened VM in scenario 1

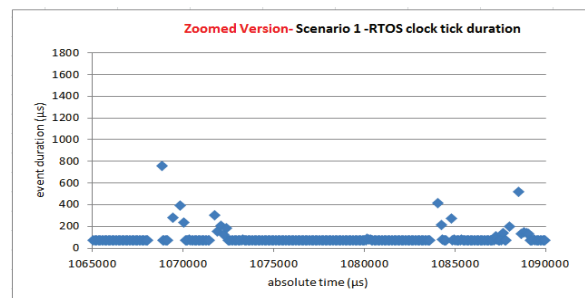


Figure 7: Zoomed version of the red sport in figure 6

Figure 6 shows that every second, the hypervisor is doing some tasks/scheduling decisions which causes the VM to be suspended/scheduled-away resulting in such high values periodically.

Note that our policy is black-box testing which makes it difficult to understand the internal behavior of an out-of-the-box product.

The emulated VM behaves exactly the same except with higher values. Table 2 is a comparison between the “clock tick processing duration” test results for both VMs, while table 3 is a comparison for the “thread switch latency” test results.

Table 2 : Comparison between the “clock-tick processing duration” test results

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	4.35 ms
Enlightened VM	1.62 ms

Table 3 : Comparison between the “Thread switch latency” test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	3.53 ms	1.24 ms
Enlightened VM	63 μ s	687 μ s

Note that the “maximum switch latency” in all the scenarios depends on the processing durations of clock tick and other interrupts that may occur in the system during the testing time.

ii. Scenario2: One-to-One

As mentioned before, Hyper-V does not support affinity. It sends the workload of a VM to the first physical CPU that is available.

In this scenario, there is only one physical CPU available, while the other three are disabled from the BIOS. There is only the UTM, together with the parent partition which is always parked on CPU-0 but idle. Therefore, UTM is also pinned to CPU-0. The aim of this scenario is to clarify if the affinity technique removes the periodic high measurements.

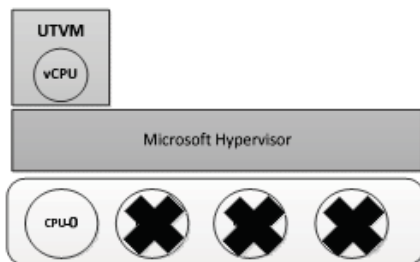


Figure 8 : One-to-One scenario

Figure 9 shows that the periodic high values are still detected.

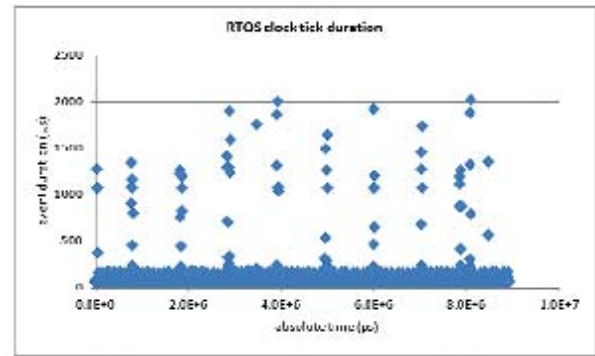


Figure 9 : Clock tick duration test for Enlightened VM in scenario 2

Tables 4 and 5 compares the results of the two tests: “Clock Tick processing duration” and “Thread switch latency”.

Table 4 : Comparison between the “clock-tick processing duration” test results

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	7.19 ms
Enlightened VM	4.08 ms

Table 5 : Comparison between the “Thread switch latency” test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	1.78 ms	1.8 ms
Enlightened VM	1.24 ms	1.47 ms

iii. Scenario3: Contention with 1 CPU-Load VM

This scenario has 2 VMs, UTM and CPU-Load VM, both running on the same physical CPU as shown in figure 10. The CPU-Load VM is running a CPU-stress program which is an infinite loop of mathematical calculations. The aim of this scenario is to explore the scheduling mechanism of the hypervisor between competing VMs.

Each of the two VMs has its “virtual machine reserve” and “virtual machine limit” attributes set to 50%.

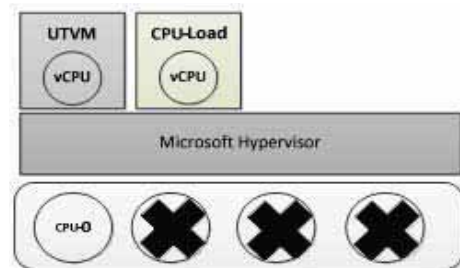


Figure 10 : Contention with 1 CPU-Load VM scenario

Tables 6 and 7 compares the results of the two tests: "Clock Tick processing duration" and "Thread switch latency".

Table 6 : Comparison between the "clock-tick processing duration" test results

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	18.56 ms
Enlightened VM	9.16 ms

Table 7 : Comparison between the "Thread switch latency" test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	5.96 ms	5.96 ms
Enlightened VM	5.7 ms	5.17 ms

iv. Scenario4: Contention with 1 Memory-Load VM

This scenario is exactly the same as scenario 3 except using Memory-Load VM instead of CPU-Load VM. This VM is running an infinite loop of memcpy() function that copies 9 MB (a value that is larger than the whole caches) from one object to another. The other goal of this scenario using such a VM is to detect the cache effects on the performance of the UTVM.

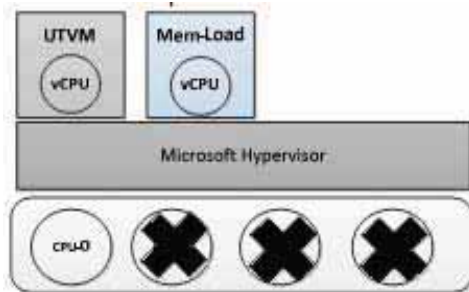


Figure 11 : Contention with 1 Memory-Load VM scenario

Tables 8 and 9 compare the results of the two tests: "Clock Tick processing duration" and "Thread switch latency".

Table 8 : Comparison between the "clock-tick processing duration" test results

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	21.3 ms
Enlightened VM	11.8 ms

Table 8 shows that measurements of this scenario are greater than the ones of the previous scenario (scenario 3) by almost 3 ms.

Table 9 : Comparison between the "Thread switch latency" test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	6.8 ms	6.8 ms
Enlightened VM	6.7 ms	5.2 ms

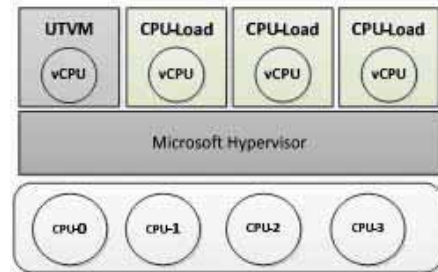


Figure 12 : All-to-All with 3 CPU-Load VMs scenario

Again, the periodic peaks are captured as shown in figure 13 below.

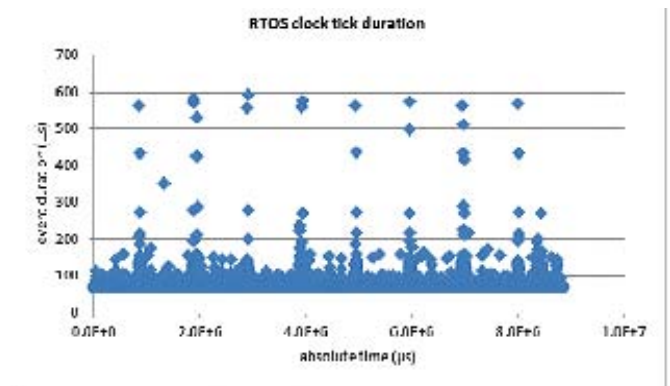


Figure 13 : Clock tick duration test for Enlightened VM in scenario 5

Tables 10 and 11 compare the results of the two tests: "Clock Tick processing duration" and "Thread switch latency".

Table 10 : Comparison between the "clock-tick processing duration" test results

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	5.21 ms
Enlightened VM	2.55 ms

Table 11 : Comparison between the "Thread switch latency" test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	76 μ s	165 μ s
Enlightened VM	36 μ s	108 μ s

vi. *Scenario6: All-to-All with 3 Memory-Load VMs*

This scenario is exactly the same as the previous scenario (scenario 5) except using Memory-Load VMs instead of CPU-Load VMs as shown in figure 14.

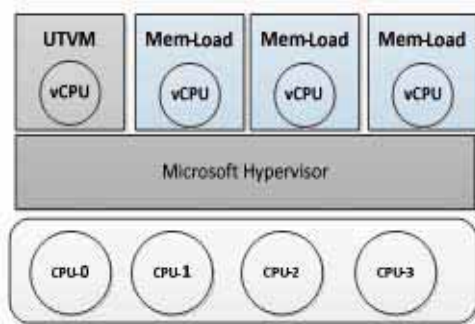


Figure 14 : All-to-All with 3 Memory-Load VMs scenario

Tables 12 and 13 compare the results of the two tests: “Clock Tick processing duration” and “Thread switch latency”.

Table 12 : Comparison between the “clock-tick processing duration” test results

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	15.18 ms
Enlightened VM	4.53 ms

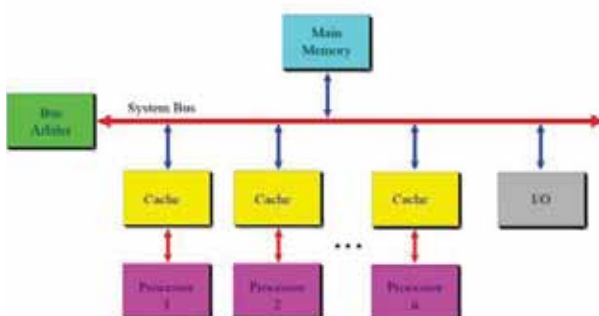
Table 13 : Comparison between the “Thread switch latency” test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	167 μ s	445 μ s
Enlightened VM	87 μ s	393 μ s

The resulting values of this scenario are around three times greater than the ones of the previous scenario (scenario 5) even though the same number of VMs is running. This difference in the results is due to the concept explained in the following section (System bus bottleneck in SMP systems).

System bus bottleneck in SMP systems.

SMP - Symmetric Multiprocessor System



The hardware platform used for this evaluation is a Symmetric Multiprocessing (SMP) system with four identical processors connected to a single shared main memory using a system bus. They have full access to all I/O devices and are treated equally.

The system memory bus or system bus can be used by only one core at a time. If two processors are executing tasks that need to use the system bus at the same time, then one of them will use the bus while the other will be blocked for some time. As the processor used has 4 cores, when all of these are running at the same time, system bus contention occurs.

Scenario 5 is not causing high overheads because the CPU stress program is quite small and fits in the core cache together with its data. Therefore, the three CPU-Loading VMs are not intensively loading the system bus which in turn will not highly affect the UTVM.

Referring back to scenario 6, the three Memory-Load VMs are intensively using the system bus. The UTVM is also running and requires the usage of system bus from time to time. Therefore, the system bus is shared most of the time between four VMs (UTVM and 3 Memory-Load VMs), which causes extra contention. Thus, the more cores in the system that are accessing the system bus simultaneously, the more contention will occur and thus the overhead increases.

To explicitly show this effect, we created another additional scenario (scenario 7 below) where only one Memory-Load VM is sharing the resources with the UTVM. The following demonstrates our observation.

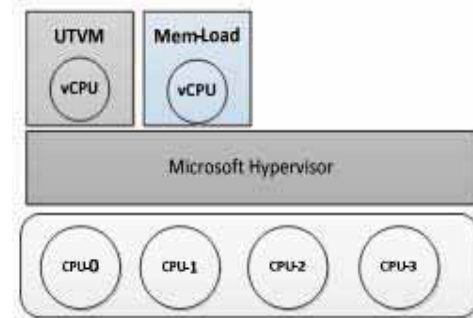
vii. *Scenario7: TWO-to-ALL with 1 Memory-Load VM*

Figure 15 : Two-to-All with 1 Memory-Load VM scenario

Tables 15 and 16 compare the results of the two tests: “Clock Tick processing duration” and “Thread switch latency”.

Clock Tick Processing Duration	
	Maximum Overhead
Bare-Machine	10 μ s
Emulated VM	5.3 ms
Enlightened VM	2.12 ms

Table 15 : Comparison between the “clock-tick processing duration” test results

Maximum Switch Latency between:		
	2 Threads	1000 Threads
Bare-Machine	11.2 μ s	11.45 μ s
Emulated VM	114 μ s	212 μ s
Enlightened VM	67 μ s	166 μ s

Table 16 : Comparison between the "Thread switch latency" test results

V. CONCLUSION

Microsoft recent Hyper-V technology is a "Microkernelized Type 1" hypervisor which leverages paravirtualization (called Enlightenment by Microsoft) in addition to the traditional hardware emulation technique. It exists in two variants: as a stand-alone product called Hyper-V Server and as an installable role in Windows Server.

There is no difference between MS Hyper-V in each of these two variants. The hypervisor is the same regardless of the installed edition.

In this paper, MS Hyper-V Server 2012 is undergoing testing. It installs a very minimal set of Windows Server components to optimize the virtualization environment. It supports Enlightened (special drivers are added to the VM to provide advanced features and performance) and hardware-emulated VMs.

This work compares the performance between the two types of VM. For this purpose, different tests and several scenarios are used. The results show that the enlightened VM performs on average twice as good as the hardware-emulated VM. This performance enhancement may increase/decrease depending on the scenario in question.

Even with this improvement, Enlightened VM performance is low compared with bare-machine (non-virtualized) performance.

A shared-memory symmetric multiprocessor hardware with four physical cores is used for conducting the tests. The results also show that the synchronous usage of all the available cores causes an intensive overload in the system bus which in turn increases latencies by a factor of 3 when compared with a system with only one active core.

REFERENCES RÉFÉRENCES REFERENCIAS

1. M. Lee, A. S. Krishnakumar, P. Krishnan, S. Nayjot and Y. Shalini, "Supporting Sofy Real-Time Tasks in the Xen Hypervisor," in the 6th ACM SIGPLAN/SIGOPS international conference on Virtual execution enviroments, 2010.
2. VMWare, "Understanding Full virtualization, Paravirtualization and hardware Assist," 2007. [Online]. Available: http://www.vmware.com/files/pdf/VMware_paravirtualization.pdf.
3. Z. H. Shah, Windows Server 2012 Hyper-V: Deploying Hyper-V Enterprise Server Virtualization Platform, Packt Publishing, 2013.
4. Virtuatopia, "An Overview of the Hyper-V Architecture," [Online]. Available: http://www.virtuatopia.com/index.php/An_Overview_of_the_HyperV_Architecture.
5. T. Abels, P. Dhawam and B. Chandrasekaran, "An overview of Xen Virtualization," [Online]. Available: <http://www.dell.com/downloads/global/power/ps3q05-20050191-abels.pdf>.
6. Microsoft, "Hyper-V Architecture," [Online]. Available: <http://msdn.microsoft.com/enus/library/cc768520%28v=bts.10%29.aspx>.
7. M. T. Blogs, "Hyper-V: Microkernelized or Monolithic," [Online]. Available: <http://blogs.technet.com/b/chenley/archive/2011/02/23/hyper-v-microkernelize-d-or-monolithic.aspx>.
8. Finn and P. Lownds, Mastering Hyper-V Deployment, Wiley Publishing Inc.
9. M. Corporation, "Windows Server 2008 Hyper-V Technical Overview," [Online]. Available: <http://download.microsoft.com>.
10. G. Knuth, "Microsoft Windows Server 2008 – Hyper-V solution overview," 2008. [Online]. Available: <http://www.brianmadden.com/blogs/gabeknuth/archive/2008/03/11/microsoft-windows-server-2008-hyper-v-solution-overview.aspx>.
11. Microsoft, "Microsoft Hyper-V Server 2012," [Online]. Available: <http://www.microsoft.com/enus/server-cloud/hyper-v-server/>.
12. "CONFIG PREEMPT RT Patch-RT wiki," [Online]. Available: https://rt.wiki.kernel.org/index.php/CONFIG_PREEMPT_RT_Patch.
13. T. B. developers, "Buildroot: Making Embedded Linux easy," [Online]. Available: <http://buildroot.uclibc.org/>.
14. B. Armstrong, "Hyper-V CPU Scheduling-Part 1," 2011. [Online]. Available: http://blogs.msdn.com/b/virtual_pc_guy/archive/2011/02/14/hyper-v-cpu-scheduling-part-1.aspx.
15. M. T. Wiki, "Hyper-V Concepts - vCPU (Virtual Processor)," [Online]. Available: <http://social.technet.microsoft.com/wiki/contents/articles/1234.hyper-v-conceptsvcpuvirtualprocessor.aspx?wa=wsignin1.0>.
16. M. H. Klein, T. Ralya, B. Pollak, R. Obenza and M. G. Harbour, A practitioner's Handbook for Real-Time Analysis, USA: Kumer Academic Publishers, 1994. ISBN 0-7923-9361-9.

GLOBAL JOURNALS INC. (US) GUIDELINES HANDBOOK 2013

WWW.GLOBALJOURNALS.ORG

FELLOWS

FELLOW OF ASSOCIATION OF RESEARCH SOCIETY IN COMPUTING (FARSC)

Global Journals Incorporate (USA) is accredited by Open Association of Research Society (OARS), U.S.A and in turn, awards “FARSC” title to individuals. The 'FARSC' title is accorded to a selected professional after the approval of the Editor-in-Chief/Editorial Board Members/Dean.



- The “FARSC” is a dignified title which is accorded to a person’s name viz. Dr. John E. Hall, Ph.D., FARSC or William Walldroff, M.S., FARSC.

FARSC accrediting is an honor. It authenticates your research activities. After recognition as FARSC, you can add 'FARSC' title with your name as you use this recognition as additional suffix to your status. This will definitely enhance and add more value and repute to your name. You may use it on your professional Counseling Materials such as CV, Resume, and Visiting Card etc.

The following benefits can be availed by you only for next three years from the date of certification:



FARSC designated members are entitled to avail a 40% discount while publishing their research papers (of a single author) with Global Journals Incorporation (USA), if the same is accepted by Editorial Board/Peer Reviewers. If you are a main author or co-author in case of multiple authors, you will be entitled to avail discount of 10%.

Once FARSC title is accorded, the Fellow is authorized to organize a symposium/seminar/conference on behalf of Global Journal Incorporation (USA). The Fellow can also participate in conference/seminar/symposium organized by another institution as representative of Global Journal. In both the cases, it is mandatory for him to discuss with us and obtain our consent.



You may join as member of the Editorial Board of Global Journals Incorporation (USA) after successful completion of three years as Fellow and as Peer Reviewer. In addition, it is also desirable that you should organize seminar/symposium/conference at least once.

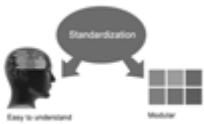
We shall provide you intimation regarding launching of e-version of journal of your stream time to time. This may be utilized in your library for the enrichment of knowledge of your students as well as it can also be helpful for the concerned faculty members.





The FARSC can go through standards of OARS. You can also play vital role if you have any suggestions so that proper amendment can take place to improve the same for the benefit of entire research community.

As FARSC, you will be given a renowned, secure and free professional email address with 100 GB of space e.g. johnhall@globaljournals.org. This will include Webmail, Spam Assassin, Email Forwarders, Auto-Responders, Email Delivery Route tracing, etc.



The FARSC will be eligible for a free application of standardization of their researches. Standardization of research will be subject to acceptability within stipulated norms as the next step after publishing in a journal. We shall depute a team of specialized research professionals who will render their services for elevating your researches to next higher level, which is worldwide open standardization.

The FARSC member can apply for grading and certification of standards of their educational and Institutional Degrees to Open Association of Research, Society U.S.A. Once you are designated as FARSC, you may send us a scanned copy of all of your credentials. OARS will verify, grade and certify them. This will be based on your academic records, quality of research papers published by you, and some more criteria. After certification of all your credentials by OARS, they will be published on your Fellow Profile link on website <https://associationofresearch.org> which will be helpful to upgrade the dignity.



The FARSC members can avail the benefits of free research podcasting in Global Research Radio with their research documents. After publishing the work, (including published elsewhere worldwide with proper authorization) you can upload your research paper with your recorded voice or you can utilize chargeable services of our professional RJs to record your paper in their voice on request.

The FARSC member also entitled to get the benefits of free research podcasting of their research documents through video clips. We can also streamline your conference videos and display your slides/ online slides and online research video clips at reasonable charges, on request.





The FARSC is eligible to earn from sales proceeds of his/her researches/reference/review Books or literature, while publishing with Global Journals. The FARSC can decide whether he/she would like to publish his/her research in a closed manner. In this case, whenever readers purchase that individual research paper for reading, maximum 60% of its profit earned as royalty by Global Journals, will be credited to his/her bank account. The entire entitled amount will be credited to his/her bank account exceeding limit of minimum fixed balance. There is no minimum time limit for collection. The FARSC member can decide its price and we can help in making the right decision.

The FARSC member is eligible to join as a paid peer reviewer at Global Journals Incorporation (USA) and can get remuneration of 15% of author fees, taken from the author of a respective paper. After reviewing 5 or more papers you can request to transfer the amount to your bank account.



MEMBER OF ASSOCIATION OF RESEARCH SOCIETY IN COMPUTING (MARSC)

The ' MARSC ' title is accorded to a selected professional after the approval of the Editor-in-Chief / Editorial Board Members/Dean.

The "MARSC" is a dignified ornament which is accorded to a person's name viz. Dr. John E. Hall, Ph.D., MARSC or William Walldroff, M.S., MARSC.



MARSC accrediting is an honor. It authenticates your research activities. After becoming MARSC, you can add 'MARSC' title with your name as you use this recognition as additional suffix to your status. This will definitely enhance and add more value and reputé to your name. You may use it on your professional Counseling Materials such as CV, Resume, Visiting Card and Name Plate etc.

The following benefits can be availed by you only for next three years from the date of certification.



MARSC designated members are entitled to avail a 25% discount while publishing their research papers (of a single author) in Global Journals Inc., if the same is accepted by our Editorial Board and Peer Reviewers. If you are a main author or co-author of a group of authors, you will get discount of 10%.

As MARSC, you will be given a renowned, secure and free professional email address with 30 GB of space e.g. johnhall@globaljournals.org. This will include Webmail, Spam Assassin, Email Forwarders, Auto-Responders, Email Delivery Route tracing, etc.





We shall provide you intimation regarding launching of e-version of journal of your stream time to time. This may be utilized in your library for the enrichment of knowledge of your students as well as it can also be helpful for the concerned faculty members.

The MARSC member can apply for approval, grading and certification of standards of their educational and Institutional Degrees to Open Association of Research, Society U.S.A.



Once you are designated as MARSC, you may send us a scanned copy of all of your credentials. OARS will verify, grade and certify them. This will be based on your academic records, quality of research papers published by you, and some more criteria.

It is mandatory to read all terms and conditions carefully.



AUXILIARY MEMBERSHIPS

Institutional Fellow of Open Association of Research Society (USA)-OARS (USA)

Global Journals Incorporation (USA) is accredited by Open Association of Research Society, U.S.A (OARS) and in turn, affiliates research institutions as “Institutional Fellow of Open Association of Research Society” (IFOARS).

The “FARSC” is a dignified title which is accorded to a person’s name viz. Dr. John E. Hall, Ph.D., FARSC or William Walldroff, M.S., FARSC.



The IFOARS institution is entitled to form a Board comprised of one Chairperson and three to five board members preferably from different streams. The Board will be recognized as “Institutional Board of Open Association of Research Society”-(IBOARS).

The Institute will be entitled to following benefits:



The IBOARS can initially review research papers of their institute and recommend them to publish with respective journal of Global Journals. It can also review the papers of other institutions after obtaining our consent. The second review will be done by peer reviewer of Global Journals Incorporation (USA). The Board is at liberty to appoint a peer reviewer with the approval of chairperson after consulting us.

The author fees of such paper may be waived off up to 40%.

The Global Journals Incorporation (USA) at its discretion can also refer double blind peer reviewed paper at their end to the board for the verification and to get recommendation for final stage of acceptance of publication.



The IBOARS can organize symposium/seminar/conference in their country on behalf of Global Journals Incorporation (USA)-OARS (USA). The terms and conditions can be discussed separately.

The Board can also play vital role by exploring and giving valuable suggestions regarding the Standards of “Open Association of Research Society, U.S.A (OARS)” so that proper amendment can take place for the benefit of entire research community. We shall provide details of particular standard only on receipt of request from the Board.



Journals Research
inducing researches

The board members can also join us as Individual Fellow with 40% discount on total fees applicable to Individual Fellow. They will be entitled to avail all the benefits as declared. Please visit Individual Fellow-sub menu of GlobalJournals.org to have more relevant details.

We shall provide you intimation regarding launching of e-version of journal of your stream time to time. This may be utilized in your library for the enrichment of knowledge of your students as well as it can also be helpful for the concerned faculty members.



After nomination of your institution as “Institutional Fellow” and constantly functioning successfully for one year, we can consider giving recognition to your institute to function as Regional/Zonal office on our behalf.

The board can also take up the additional allied activities for betterment after our consultation.

The following entitlements are applicable to individual Fellows:

Open Association of Research Society, U.S.A (OARS) By-laws states that an individual Fellow may use the designations as applicable, or the corresponding initials. The Credentials of individual Fellow and Associate designations signify that the individual has gained knowledge of the fundamental concepts. One is magnanimous and proficient in an expertise course covering the professional code of conduct, and follows recognized standards of practice.



Open Association of Research Society (US)/ Global Journals Incorporation (USA), as described in Corporate Statements, are educational, research publishing and professional membership organizations. Achieving our individual Fellow or Associate status is based mainly on meeting stated educational research requirements.

Disbursement of 40% Royalty earned through Global Journals : Researcher = 50%, Peer Reviewer = 37.50%, Institution = 12.50% E.g. Out of 40%, the 20% benefit should be passed on to researcher, 15 % benefit towards remuneration should be given to a reviewer and remaining 5% is to be retained by the institution.



We shall provide print version of 12 issues of any three journals [as per your requirement] out of our 38 journals worth \$ 2376 USD.

Other:

The individual Fellow and Associate designations accredited by Open Association of Research Society (US) credentials signify guarantees following achievements:

- The professional accredited with Fellow honor, is entitled to various benefits viz. name, fame, honor, regular flow of income, secured bright future, social status etc.



- In addition to above, if one is single author, then entitled to 40% discount on publishing research paper and can get 10% discount if one is co-author or main author among group of authors.
- The Fellow can organize symposium/seminar/conference on behalf of Global Journals Incorporation (USA) and he/she can also attend the same organized by other institutes on behalf of Global Journals.
- The Fellow can become member of Editorial Board Member after completing 3yrs.
- The Fellow can earn 60% of sales proceeds from the sale of reference/review books/literature/publishing of research paper.
- Fellow can also join as paid peer reviewer and earn 15% remuneration of author charges and can also get an opportunity to join as member of the Editorial Board of Global Journals Incorporation (USA)
- • This individual has learned the basic methods of applying those concepts and techniques to common challenging situations. This individual has further demonstrated an in-depth understanding of the application of suitable techniques to a particular area of research practice.

Note :

//

- In future, if the board feels the necessity to change any board member, the same can be done with the consent of the chairperson along with anyone board member without our approval.
- In case, the chairperson needs to be replaced then consent of 2/3rd board members are required and they are also required to jointly pass the resolution copy of which should be sent to us. In such case, it will be compulsory to obtain our approval before replacement.
- In case of “Difference of Opinion [if any]” among the Board members, our decision will be final and binding to everyone.

//

PROCESS OF SUBMISSION OF RESEARCH PAPER

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

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

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

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

(II) Choose corresponding Journal.

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

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

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

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



PREFERRED AUTHOR GUIDELINES

MANUSCRIPT STYLE INSTRUCTION (Must be strictly followed)

Page Size: 8.27" X 11"

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

You can use your own standard format also.

Author Guidelines:

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

1. GENERAL

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

Scope

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

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

2. ETHICAL GUIDELINES

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

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

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

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

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

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

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

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

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

Please mention proper reference and appropriate acknowledgements wherever expected.

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

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

3. SUBMISSION OF MANUSCRIPTS

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

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



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

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

4. MANUSCRIPT'S CATEGORY

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

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

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

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

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

5. STRUCTURE AND FORMAT OF MANUSCRIPT

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

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

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

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



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

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

Format

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

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

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

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

Structure

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

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

Abstract, used in Original Papers and Reviews:

Optimizing Abstract for Search Engines

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

Key Words

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

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

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

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



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

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

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

Acknowledgements: Please make these as concise as possible.

References

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

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

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

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

Tables, Figures and Figure Legends

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

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

Preparation of Electronic Figures for Publication

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

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

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



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

6. AFTER ACCEPTANCE

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

6.1 Proof Corrections

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

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

(Free of charge) from the following website:

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

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

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

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

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

6.3 Author Services

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

6.4 Author Material Archive Policy

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

6.5 Offprint and Extra Copies

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

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



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

TECHNIQUES FOR WRITING A GOOD QUALITY RESEARCH PAPER:

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

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

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

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

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

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

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

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

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

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

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



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

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

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

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

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

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

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

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

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

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

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

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

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

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

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



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

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

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

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

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

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

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

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

INFORMAL GUIDELINES OF RESEARCH PAPER WRITING

Key points to remember:

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

Final Points:

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

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



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

General style:

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

To make a paper clear

- Adhere to recommended page limits

Mistakes to evade

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

In every sections of your document

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

Title Page:

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



Abstract:

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

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

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

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

Approach:

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

Introduction:

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

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

Approach:

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



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

Procedures (Methods and Materials):

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

Materials:

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

Methods:

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

Approach:

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

What to keep away from

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

Results:

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

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



Content

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

What to stay away from

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

Approach

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

Figures and tables

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

Discussion:

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

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

Approach:

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



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

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

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

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



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

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

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



INDEX

B

Baremachine · 26
Bayesian · 2, 3

C

Collaboratively · 17
Corasick · 15

E

Endorsement · 5, 6, 7
Enlightenments · 23, 25

H

Hardwareassisted · 23

L

Linguistical · 2

M

Microkernelized · 23

P

Parallelism · 15, 18, 21

R

Repositories · 10

U

Unwillingness · 1

V

Voluntarily · 27

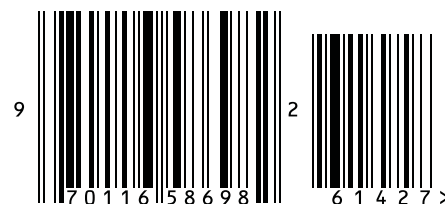


save our planet



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