

GLOBAL JOURNAL

OF COMPUTER SCIENCE AND TECHNOLOGY: A

Hardware & Computation

Digital Logic Design

Support Vector Machine

Highlights

Parallel String Matching

Hardware Synthesis of Chip

Discovering Thoughts, Inventing Future

VOLUME 13

ISSUE 1

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 1 (VER. 1.0)

OPEN ASSOCIATION OF RESEARCH SOCIETY

© Global Journal of Computer Science and Technology. 2013.

All rights reserved.

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

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

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

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

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

Engage with the contents herein at your own risk.

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

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

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

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

Global Journals Inc.

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

Sponsors: *Open Association of Research Society*
Open Scientific Standards

Publisher's Headquarters office

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

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

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

Offset Typesetting

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

Packaging & Continental Dispatching

Global Journals, India

Find a correspondence nodal officer near you

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

eContacts

Press Inquiries: press@globaljournals.org

Investor Inquiries: investers@globaljournals.org

Technical Support: technology@globaljournals.org

Media & Releases: media@globaljournals.org

Pricing (Including by Air Parcel Charges):

For Authors:

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

Yearly Subscription (Personal & Institutional):

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

EDITORIAL BOARD MEMBERS (HON.)

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

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

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

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

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

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

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

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

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

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

Dr. Bart Lambrecht

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

Dr. Carlos García Pont

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

Dr. Fotini Labropulu

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

Dr. Lynn Lim

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

Dr. Mihaly Mezei

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

Dr. Söhnke M. Bartram

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

Dr. Miguel Angel Ariño

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

Philip G. Moscoso

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

Dr. Sanjay Dixit, M.D.

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

Dr. Han-Xiang Deng

MD., Ph.D
Associate Professor and Research
Department Division of Neuromuscular
Medicine
David R. Davies 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. Comparative Analysis of Spatio and Viterbi Encoding and Decoding Techniques in Hardware Description Language. *1-9*
 - 2. Modification of Support Vector Machine for Microarray Data Analysis. *11-15*
 - 3. Hardware Synthesis of Chip Enhancement Transformations in Hardware Description Language Environment. *17-26*
 - 4. Parallel String Matching with Multi Core Processors-A Comparative Study for Gene Sequences. *27-41*
 - 5. Analysis of Parallel Boyer-Moore String Search Algorithm. *43-47*
 - 6. Software Defined 8, 16 & 24 Bit Digital Logic Design by One Microcontroller. *49-54*
-
- vii. Auxiliary Memberships
 - viii. Process of Submission of Research Paper
 - ix. Preferred Author Guidelines
 - x. Index



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 1 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Comparative Analysis of Spatio and Viterbi Encoding and Decoding Techniques in Hardware Description Language

By Pooja Nagwal, Adesh Kumar & Dharendra Singh Gangwar

University of Petroleum and Energy Studies Dehradun, India

Abstract - The paper focuses on the design and synthesis of hardware chip for Spatio and Viterbi encoding and decoding techniques. Both techniques are used for digital data encoding and decoding in transmitter and receiver respectively. These techniques are used for error control coding found in convolution codes. Spatio coding is also used to eliminate crosstalk among interconnect wires, thereby reducing delay. The encoded data in packet form may be of 'N' bits. Data is decoded at different clock pulses at which it is encoded. A comparative analysis is done for hardware parameter, timing parameters and device utilization. Design is implemented in Xilinx 14.2 VHDL software, and functional simulation was carried out in Modelsim 10.1 b, student edition. Hardware parameters such as size cost and timings are extracted from the design code.

Keywords : *field programmable gate array (FPGA), register transfer level (RTL), very high speed integrated circuit hardware description language (VHDL), very large scale of integration (VLSI).*

GJCST-A Classification : *B.m*



COMPARATIVE ANALYSIS OF SPATIO AND VITERBI ENCODING AND DECODING TECHNIQUES IN HARDWARE DESCRIPTION LANGUAGE

Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

Comparative Analysis of Spatio and Viterbi Encoding and Decoding Techniques in Hardware Description Language

Pooja Nagwal ^α, Adesh Kumar ^σ & Dharendra Singh Gangwar ^ρ

Abstract - The paper focuses on the design and synthesis of hardware chip for Spatio and Viterbi encoding and decoding techniques. Both techniques are used for digital data encoding and decoding in transmitter and receiver respectively. These techniques are used for error control coding found in convolution codes. Spatio coding is also used to eliminate crosstalk among interconnect wires, thereby reducing delay. The encoded data in packet form may be of 'N' bits. Data is decoded at different clock pluses at which it is encoded. A comparative analysis is done for hardware parameter, timing parameters and device utilization. Design is implemented in Xilinx 14.2 VHDL software, and functional simulation was carried out in Modelsim 10.1 b, student edition. Hardware parameters such as size cost and timings are extracted from the design code.

Keywords : field programmable gate array (FPGA), register transfer level (RTL), very high speed integrated circuit hardware description language (VHDL), very large scale of integration (VLSI).

1. INTRODUCTION

The Viterbi Algorithm [1] [3] [16] is used widely for estimation and detection problems in digital communication and signal processing. It is used to detect signals in communication channels with memory [3], and to decode sequential error control codes [16] that are used to improve the performance of digital communication systems. During the transmission or storage process, the digital data may get corrupted due to noise. Channel coding [1] is a method to encode the data in a manner that it can be recovered even if it gets influenced by noise. Channel Coding involves adding redundant bits [3] to the data so that when it gets corrupted due to noise, the data can still be recovered through the redundancy [1] present in it. Block codes and convolution codes [16] are two major forms of Channel coding. The block codes [16] transform a block of k symbols into a block of n symbols called code word, where $n > k$. Since the output n -symbol code word depends only upon the corresponding k -symbol

input code word, the encoder is memory less and can be implemented in a combinational logic. On the other hand, a Viterbi encoder [3] [16] not only depends on the corresponding k -symbol input block but also on m previous input blocks. Viterbi encoders are implemented in sequential logic because they are associated with memory element. Keeping in mind the essentials of communication channels in wireless systems, reliable data communication, fast as well as accurate is the main requirement and Viterbi coding helps us in achieving the same. The Viterbi algorithm [1] applies the maximum-likelihood path [3] [16] method for error detection. The most common metric used is the branch distance metric [1]. This is basically the dot product between the received codeword and the allowable codeword [16]. These metrics are cumulative so that the path with the largest total metric is the final desired output. The selection of survivor path basically determines the whole of the Viterbi algorithm and ensures that the algorithm completes with the maximum likelihood path [1]. The algorithm ends and is completed when all of the nodes in the trellis have been labeled and their entering survivor paths have been determined. The design space for VLSI implementation of Viterbi decoders is large, involving choices of throughput, latency, area, and power. Even for a fixed set of parameters like constraint length, encoder polynomials [3] and trace-back depth [1], the task of designing a Viterbi decoder is quite complicated and requires significant effort. Sometimes, due to incomplete design space exploration or incorrect analysis, a suboptimal [3] design is selected

In onchip interconnects [2] [4], there is propagation delay [5] due to resistance and inter wire interconnects and gates and some others sources are such as alpha particles, electromagnetic interference [2] and power grid fluctuation [8]. Various techniques are used to minimize the delay and also various error detection and correction scheme [6]. In this paper, a Spatio- temporal bus encoding scheme [2] [4] [5] [9] is proposed which are reduced the delay due to its optimized hardware parameters and implementation. Experimental results of this scheme is show that this scheme perform better and have advantages of error detection, also in this paper we compare this scheme with Viterbi encoding scheme.

Author ^α : M.Tech Scholar, VLSI Design, Uttarakhand Technical University, Dehradun India. E-mail : nagwal.pooja@gmail.com

Author ^σ : Assistant Professor, Department of Electrical, Electronics & Instrumentation Engineering, University of Petroleum & Energy Studies, Dehradun India. E-mail : adeshmanav@gmail.com

Author ^ρ : Assistant Professor, Faculty of Technology, Uttarakhand Technical University, Dehradun India. E-mail : dsgangwar@gmail.com

The rest of this paper is organized as follow. The algorithm of Spatio temporal scheme is proposed in section II. The algorithm of Viterbi encoding scheme is presented in section III. Section IV the presents the simulation results, RTL views and discussion part. Comparative analysis of Synthesis report and timing parameters are listed in section V.

II. SPATIO TEMPORAL BUS ENCODING & DECODING

Spatio encoding [12] is applicable for arbitrary bus encoding. This techniques are proposed for 8 bits, 16 bits or 'N' bits data. Spatio temporal bus encoding scheme is proposed for eliminates the crosstalk classes [11] [13] for large energy consumption and delay of the buses [12]. This scheme is also designed to have built-in error detection [10] with very less circuit overhead. The architectures for the encoder and decoder circuit of the Spatio temporal encoding scheme are given in

figure 1(a) and 1(b). The architecture of the Spatio encoder is proposed for scheme an 8 bits data bus. In the encoder which have data d_i and previous encoded data E_{t-1} . There are two multiplexers [11] which have 3 common inputs from data d_i and two XOR gates. It's output is E_t and E_{t+1} . The data sent on the bus at time instance $t-1$ is stored in a register of 9 bits. It is denoted by E_{t-1} . The present data is stored in register which are denoted by d_i . First multiplexer (2×1) has two inputs which are common may be any one bit of data $d_i(1)$, $d_i(2)$ and $d_i(3)$. The selection line of this multiplexer is directly configured as XOR output [12] of $d_i(2)$ and $E_{t-1}(2)$. The output of this multiplexer is common input [13] [14] for $E_t(1)$, $E_t(2)$, $E_x(3)$, $E_x(4)$. Another multiplexer also has common inputs lines which are $d_i(5)$, $d_i(6)$ and $d_i(7)$ followed by XORED selection line of $d_i(6)$ and $E_{t-1}(8)$. The decoding method for the proposed scheme is similar [15].

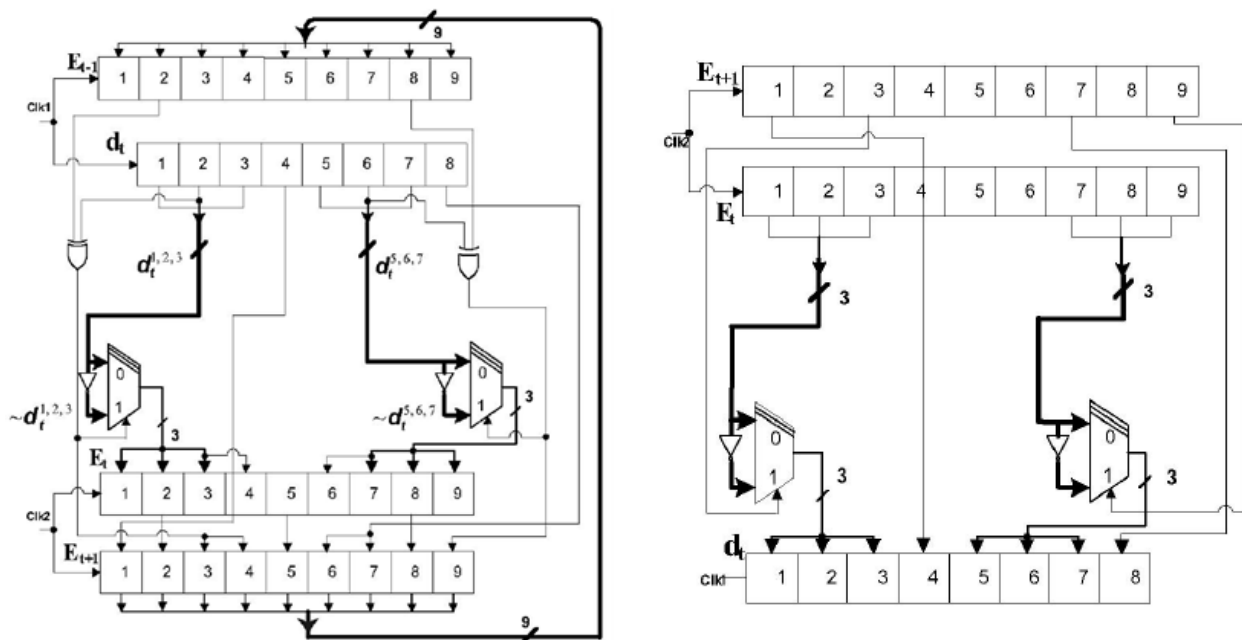


Figure 1: (a) Spatio-Encoder & (b) Spatio-Decoder

In decoding scheme [2] [12] [15] the original data d_i is reconstructed from E_t and E_{t+1} . Where E_{t+1} and E_t are the input to the decoder and d_i is obtained by the decoding algorithm. The first three bits of $d_i(1)$, $d_i(2)$ and $d_i(3)$ are common output of first multiplexer. This multiplexer accepts common input of $E_t(1)$, $E_t(2)$ and $E_t(3)$. $E_{t+1}(3)$ is the selection logic of this multiplexer. First bit of $E_{t+1}(1)$ is directly configured with $d_i(4)$. Another multiplexer accepts common inputs of $E_t(7)$, $E_t(8)$ and $E_t(9)$. Selection logic of this multiplexer is $E_{t+1}(9)$. The output of this multiplexer gives the values of $d_i(5)$, $d_i(6)$ and $d_i(7)$. $d_i(8)$ is directly configured with $E_{t+1}(7)$. For an example, consider an 8-bit data bus [2] for which encoded data will be of 9-bit length. Let the data be already available on the bus (9 bit) [2] [12] as

E_t : 101 100 010. Data to be sent on the bus d_i : 0101 101. Before coding (E_{t-1}, d_i): $\downarrow \uparrow \downarrow - \uparrow - \uparrow -$. Output of the encoder E_t : 101 100 010. Output of the encoder [2] E_{t+1} : 101 101 111. After coding (E_{t-1}, E_t): $- - - - - \uparrow \uparrow - \uparrow$. Here transaction 0-1, 1-0 and 1-1 are represented by $\uparrow, \downarrow$ and $-$.

III. VITERBI ENCODING & DECODING

Conventional codes are helping to analysis the Viterbi algorithm [3]. Viterbi algorithm are supported by two steps, the initial step is to select the trellis from the bits that are achieved at the input at the receiver. A simple trellis [1] show with 4 stage points for transmission, each state is represented with a dot and

the state transition is shown as edge of branch. Each branch is known as the branch matrix. The use of trellis structure [16] is to find the coded sequence in transmission signal. Considering a Viterbi encoder as shown in figure 2, with three modulo-2 adders, which accepts the 4 bits data stream.

Let, K = No of the shift registers = 3
 V = No. of bits in the code blocks = 3
 L = length of input data stream = 4.

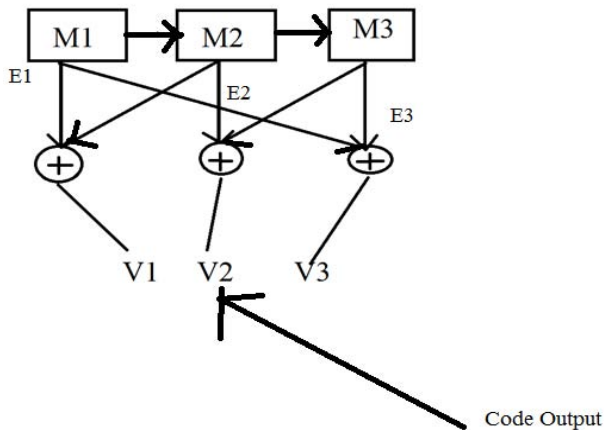


Figure 2 : Viterbi Encoder [16]

The Encoded vector $V = [V_1 \ V_2 \ V_3]$ in which the values of V_1 , V_2 and V_3 of the adders are $V_1 = E_1 \text{ XOR } E_2$, $V_2 = E_2 \text{ XOR } E_3$, $V_3 = E_1 \text{ XOR } E_3$. Initially, it is assumed that the shift register is clear means the contents of $[M_1, M_2, M_3 = 000]$. Let the 4-bit data is 1101. The data stream is entered in the shift register from MSB. MSB bit of input data stream is entered into shift register and there is one bit right shift. Thus at first bit interval is $E_1=0, E_2=1, E_3=0$. So the vector corresponding to first bit is determined by the calculation,

$$V_1 = 1 \oplus 0 = 1, V_2 = 0 \oplus 0 = 0, V_3 = 1 \oplus 0 = 1$$

Hence the value of first encoded vector is 101. Similarly second bit from MSB is entered in shift register, second bit interval $E_1=1, E_2=1, E_3=0$. Vector corresponding to second bit is

$$V_1 = 1 \oplus 1 = 0, V_2 = 1 \oplus 0 = 1, V_3 = 1 \oplus 0 = 1$$

Hence the value of second encoded vector is 101. In the same manner encoded vectors [1][16] for

other bit intervals can be found. The register will reset at seventh bit interval because maximum condition [16] of reset is $(L+K= 4+3= 7)$. The output at each bit interval consists of V bits. Thus for each message there are $(L+K)$ encoded vectors in the output code word. In the given table the coded output bit stream for all input data stream for encoder.

Table 1 : Encoded data vector for 4 bit data stream [16]

Input data stream	Coded output bit stream						
0000	000	000	000	000	000	000	000
0001	000	000	000	101	110	011	000
0010	000	000	101	110	011	000	000
0011	000	000	101	011	101	011	000
0100	000	101	110	011	000	000	000
0101	000	101	110	110	110	011	000
0110	000	101	011	101	011	000	000
0111	000	101	011	000	101	011	000
1000	101	110	011	000	000	000	000
1001	101	110	011	101	110	011	000
1010	101	110	110	110	011	000	000
1011	101	110	110	011	101	011	000
1100	101	011	101	011	000	000	000
1101	101	011	101	110	110	011	000
1110	101	011	000	101	011	000	000
1111	101	011	000	000	101	011	000

Similarly, the encoding of 8 bit can understand. Considering and 8 bit input stream 10111010. Initially shift register contents are 000. First bit of MSB is entered in shift register then $E_1=1, E_2=0, E_3=0$ the first encoded vector calculation

$$V_1 = 1 \oplus 0 = 1, V_2 = 0 \oplus 0 = 0, V_3 = 0 \oplus 1 = 1$$

Hence the value of first encoded vector is 101. Similarly second bit is entered from MSB, the contents of shift register will be $E_1=0, E_2=1, E_3=0$, the conceded vector

$$V_1 = 0 \oplus 1 = 1, V_2 = 1 \oplus 0 = 1, V_3 = 0 \oplus 0 = 0$$

Hence the value of second encoded vector is 110. The register will reset at seventh bit interval because maximum condition of reset is $(L+K= 8+3= 11)$. In the same manner encoded vectors for other bit intervals can be found and the encoded vectors are 11. Table 2 lists the values of encoded vectors for 8 bits input data stream.

Table 2 : Encoded data vector for 8 bits data stream

Input data stream	Coded output bit stream									
0000000	000	000	000	000	000	000	000	000	000	000
0000001	000	000	000	000	000	000	101	110	011	000
:	:	:	:	:	:	:	:	:	:	:
1111110	101	011	000	000	000	000	101	011	000	000
1111111	101	011	000	000	000	000	000	101	011	000

The Viterbi algorithm applies the maximum-likelihood principle [3]. The most common metric used is the Hamming distance metric [1]. This is just the dot product between the received codeword and the allowable codeword. These metrics are cumulative so that the path with the largest total metric is the final

winner. The selection of survivor path is the main feature of the Viterbi algorithm and ensures that the algorithm terminates with the maximum likelihood path. The algorithm terminates when all of the nodes in the trellis have been labeled and their entering survivor paths are determined [16].

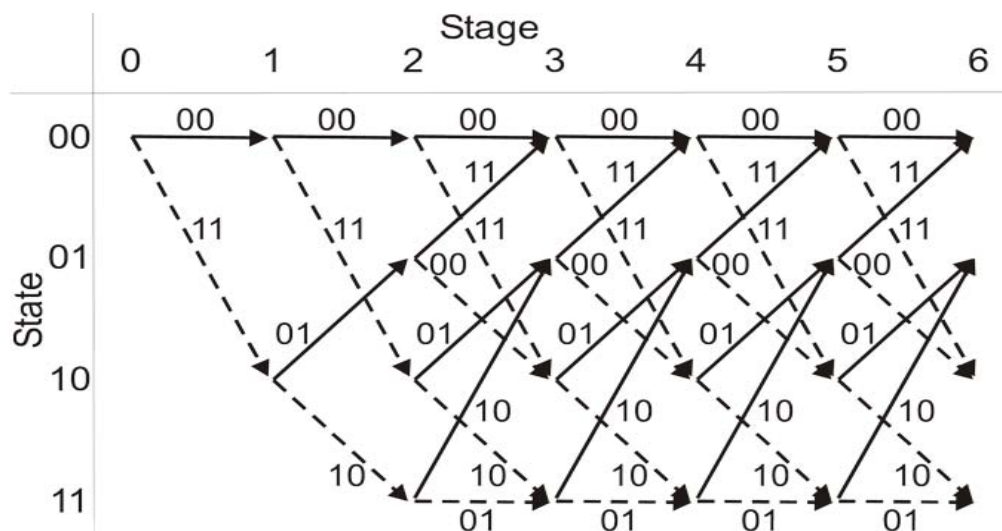


Figure 3: Trellis diagram of decoding logic [1]

IV. SIMULATION RESULT & DISCUSSION

The snapshot shown in figure 4 (a), (b) and 5(1),(b) is taken from the modelsim 10.1b software which shows the 8-bit data encoding and decoding

using Spatio and Viterbi Algorithms respectively. Register Transfer Logic (RTL) representations of both schemes are shown in the figure 6 and 7 respectively. Table 3 explains the role of pins and their functions.

Table 3 : Function Description of Pins

Pins	Functional Description
clk	Signal produce to Clock signal (1 bit of std_logic)
Reset	used for synchronization of the components by using clk (1 bit of std_logic)
Tx	9 bit encoded data by Spatio encoder at time t
Tx_Minus	9 bit encoded data by Spatio encoder at time t-1
Tx_plus	9 bit encoded data by Spatio encoder at time t+1
dt	Input data of Spatio Encoder and output of Spatio Decoder (8 bits of std_logic_vector)
Data_stream	Input data of 8 bit for Viterbi Encoder (8 bits of std_logic_vector)
Encoded_vector_data	Array of Encoded data (0 to 10) for 8 bit data_stream(8 bits of std_logic_vector)
Decoded_data_stream	Decoded output of Viterbi decoder (8 bits of std_logic_vector)

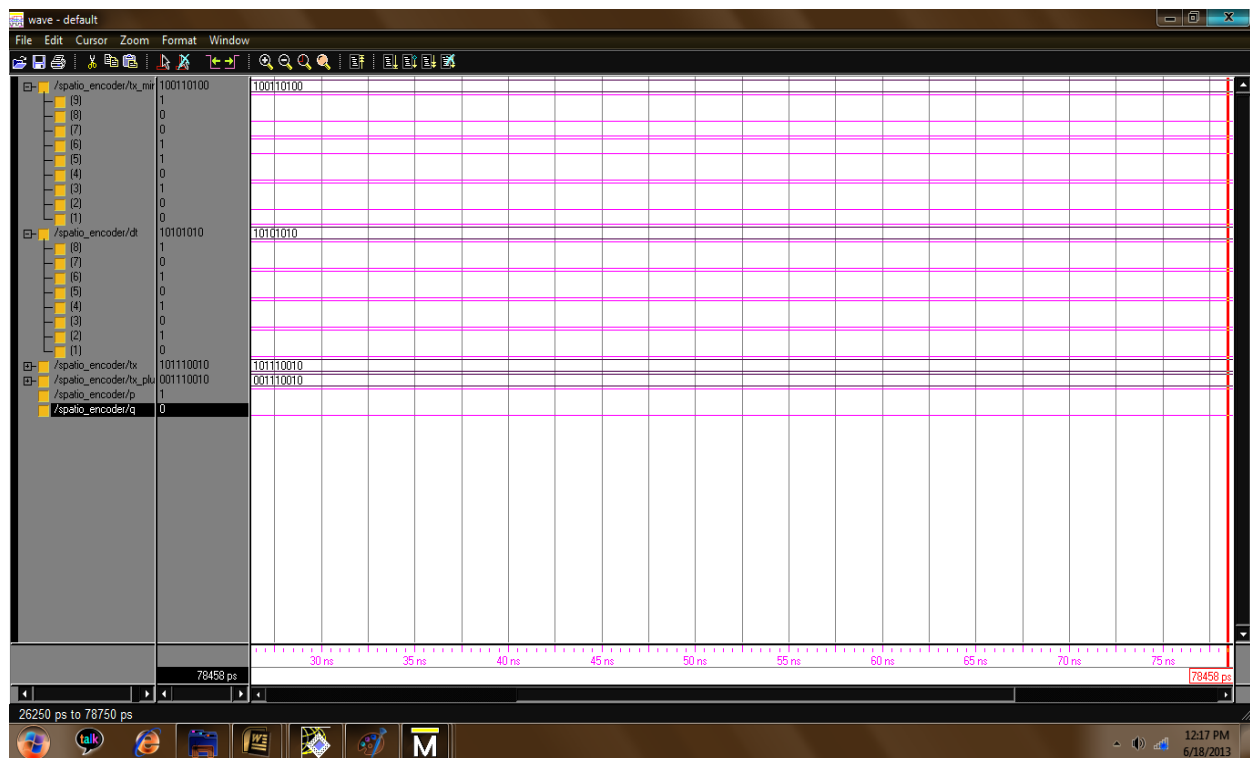


Figure 4 (a) : Modelsim waveform of 8-bits Spatio-temporal bus encoder

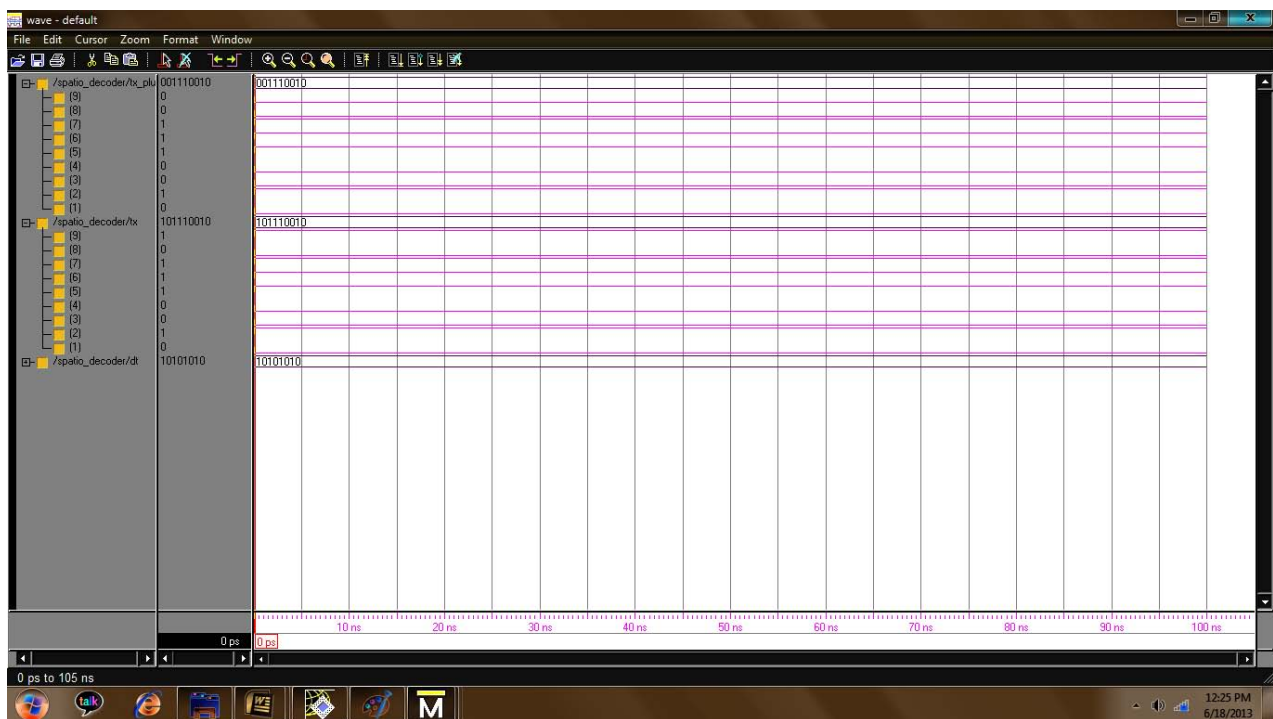


Figure 4 (b) : Modelsim waveform of 8-bits Spatio-temporal bus decoder

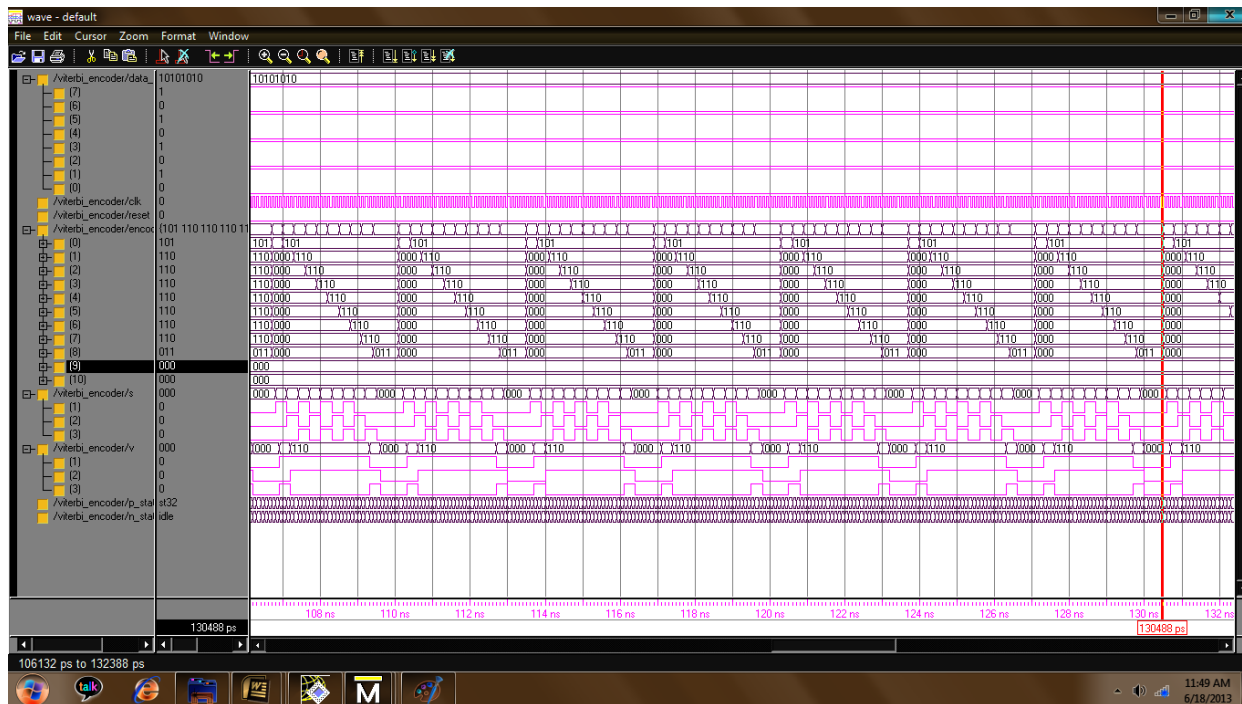


Figure 5 (a) : Modelsim waveform of 8-bits viterbi bus encoder

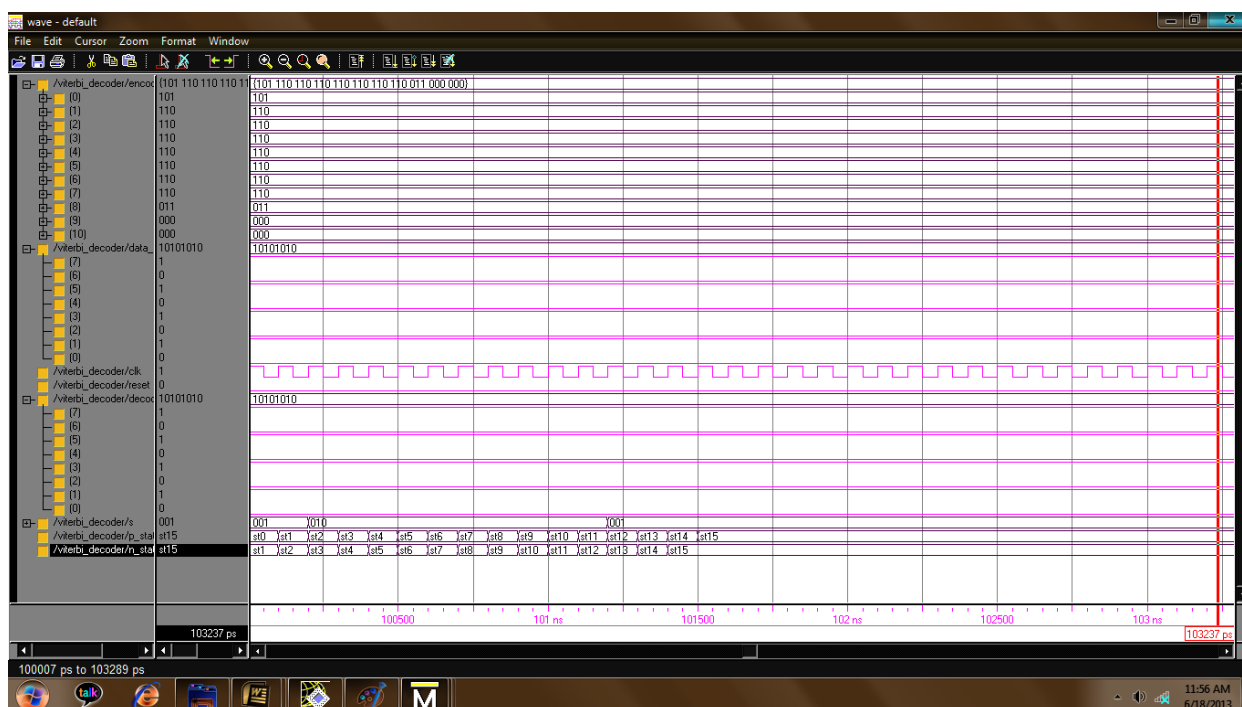


Figure 5 (b) : Modelsim waveform of 8-bit viterbi bus decoder

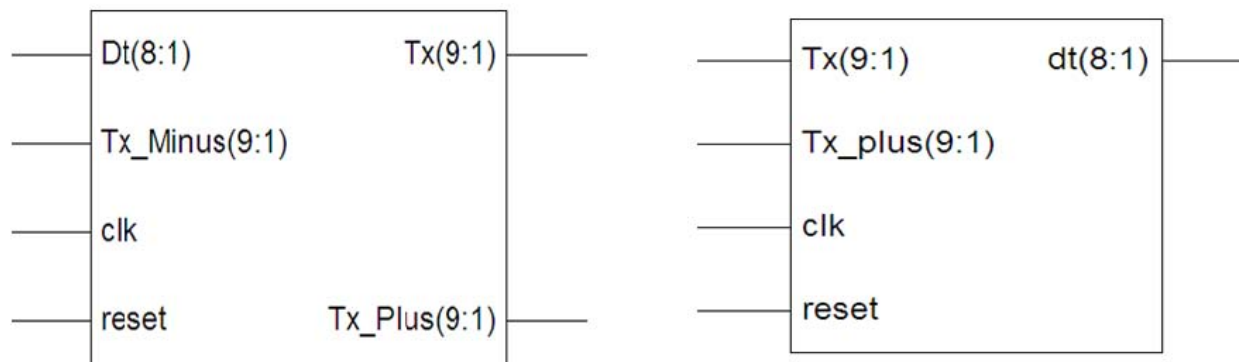


Figure 6 : (a) Spatio Encoder & (b) Decoder

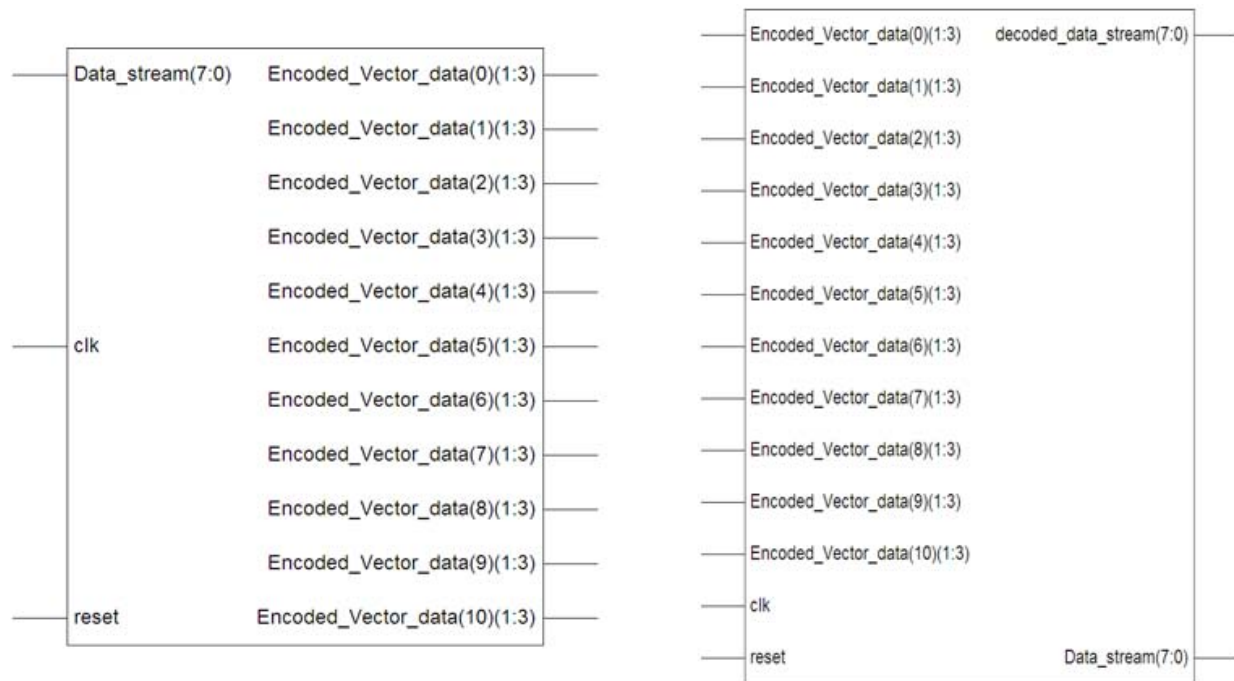


Figure 7 : (a) Viterbi Encoder & (b) Decoder

V. COMPARATIVE ANALYSIS

Comparative analysis of Spatio and Viterbi encoding and decoding Schemes can be done by the timing parameters and device utilization summary extracted from Xilinx. Device utilization summary is the report of used device hardware in the implementation of the chip such as RAM, ROM, slices, flip flops etc. Synthesis report shows the complete details of device utilization as total memory utilization. If synthesis report does not have the optimized hardware, further chip development can be done in the Xilinx ISE design software. The device targeted for synthesis on SPARTEN-3E FPGA .Timing parameters are synchronized with the clock signal. Timing details provides the information of net delay, minimum period, minimum input arrival time before clock and maximum output required time after clock. Table 4 compares the

hardware utilization for Spatio and Viterbi Encoders. Table 5 compares the hardware utilization for Spatio and Viterbi decoders. Table 6 shows the timing parameter of Spatio and Viterbi Encoders and decoders.

Selected Device : xc3s250e-5pq208

Table 4 : Hardware Utilization of Viterbi and Spatio Encoders

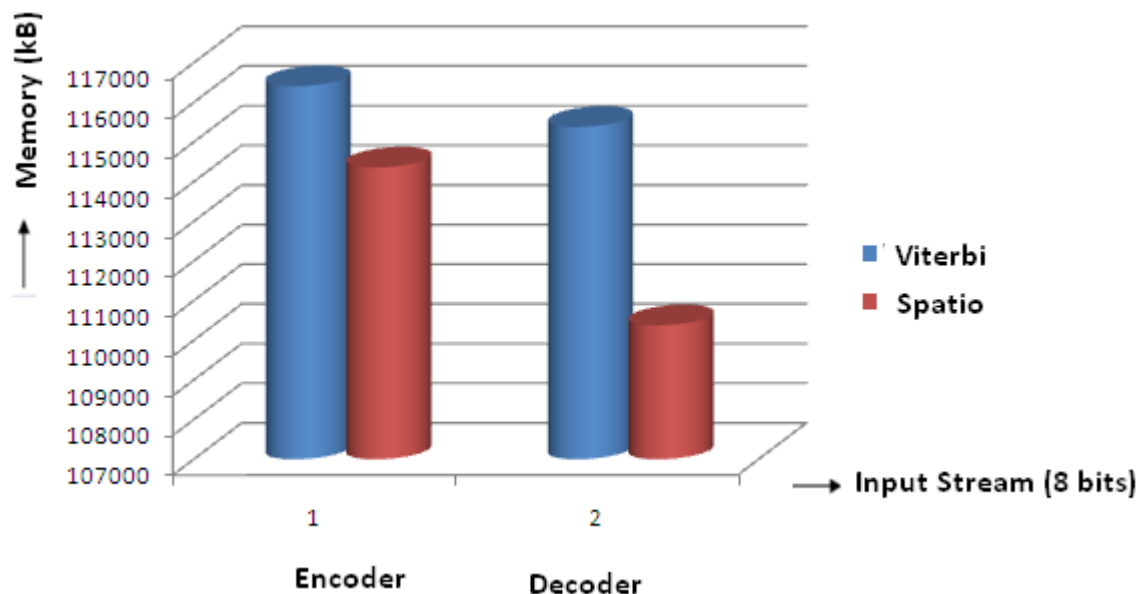
Device part	Viterbi Encoder			Spatio Encoder		
	Used	Available	Utilization	Used	Available	Utilization
Number of Slices	60	2448	2%	1	2448	0 %
Number of Slice Flip Flops	23	4896	0 %	2	4896	0 %
Number of 4 input LUTs	105	4896	2 %	50	4896	1%
Number of bonded IOBs	43	158	27 %	46	158	29 %
Number of GCLKs	1	24	4 %	1	24	4 %

Table 5 : Hardware Utilization of Viterbi and Spatio Decoders

Device part	Viterbi Decoder			Spatio Decoder		
	Used	Available	Utilization	Used	Available	Utilization
Number of Slices	34	2448	1%	26	2448	1%
Number of Slice Flip Flops	43	4896	0 %	33	4896	0 %
Number of 4 input LUTs	6	4896	0 %	4	4896	0 %
Number of bonded IOBs	18	158	11 %	16	258	10 %
Number of GCLKs	1	24	4 %	1	24	4 %

Table 6 : Timing Parameters of Viterbi and Spatio Encoders and Decoders

Parameter	Utilization			
	Viterbi Encoder	Viterbi Decoder	Spatio Encoder	Spatio Decoder
Minimum Period	3.047ns	2.058 ns	1.578 ns	1.567 ns
Minimum input arrival time before clock	8.515ns	4.053 ns	2.056 ns	2.023 ns
Maximum output required time after clock	4.179ns	4.179 ns	3.0567 ns	3.067 ns
Maximum combinational path delay	12.014 ns	10.057 ns	6.232ns	5.797 ns
Maximum Frequency	325.165 MHz	325.53 MHz	325 .27 MHz	325 .10 MHz
Memory Utilization	116408 kB	115384 kB	114360 kB	110380 kB

*Figure 8* : Memory Utilization Graph

Device utilization summary shows that there is very less difference in hardware utilization in both encoding and decoding scheme. There is 2 % difference in the Number of bonded IOBs, 2 %

difference in number of slices, 1% difference in Number of 4 input LUTs with respect to Viterbi encoder and Spatio encoder. Similarly, there is 1 % less number of bounded I/Os for Spatio decoder than Viterbi decoder.

The memory utilization graph is shown in figure 8 which shows 1.75 % less memory for Spatio encoder 4.33 % less memory for Spatio decoder. There is a reduction of 48 % in minimum period, 76 % in minimum input arrival time before clock, 27 % Maximum output required time after clock and 48 % in combinational path delay in Spatio encoder in comparison to Viterbi encoder. Similarly, There is a reduction of 24 % in minimum period, 50 % in minimum input arrival time before clock, 26 % Maximum output required time after clock and 42 % in combinational path delay in Spatio decoder in comparison to Viterbi decoder.

VI. CONCLUSION

The hardware chip for Viterbi encoder and decoder, Spatio encoder and decoder is implemented in Xilinx 14.2 and functionally checked in Modelsim 10.1b software. A comparative analysis is done with respect to hardware and timing parameters. In the chip implementation, it is analyzed that Spatio encoding and decoding is having less delay in comparison to Viterbi encoding and decoding. Memory utilization is less which is 1.75 % for Spatio encoder 4.33 % for Spatio decoder in comparison to Viterbi encoder and decoder. But there is a reduction of 48 % in combinational path delay in Spatio encoder and 42 % in Spatio decoder in comparison to Viterbi encoder and decoder. Hence Spatio encoding and decoding scheme is faster in comparison to Viterbi encoding and decoding.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Anubhuti Khare, Manish Saxena, Jagdish Patel "FPGA Based Efficient Implementation of Viterbi Decoder" International Journal of Engineering and Advanced Technology, October 2011.
2. C. Raghunandan, K. S. Sainarayanan, M. B. Srinivas "Encoding with repeater Insertion for Minimizing delay in VLSI interconnects" 2006 IEEE.
3. Chitra. M, A. R Ashwath, Roopa. M "design and Implementation of Viterbi Encoder and Decoder Using FPGA" International Journal of Engineering and Advanced Technology june 2012.
4. George Kornaros "Temporal Coding Schemes for Energy Efficient Data Transmission in Systems-on-chip".
5. J.V.R. Ravindra, Navya Chittavu, M.B. Srinivas "Energy Efficient Spatial coding Technique for Low power VLSI Application" the 6th International Workshop on System on Chip for Real Time Application 2006 IEEE.
6. J. V. R. Ravindra, K. S. Sainarayanan, M. B. Srinivas "An efficient power reduction technique for low power data I/O for military application" 2005.
7. K. S. Sainarayanan C. Raghunandan M. B. Srinivas "Delay and Power Minimization in VLSI Interconnects with Spatio-temporal Bus-Encoding Scheme" IEEE Computer Society Annual Symposium on VLSI 2007.
8. K. S. Sainarayanan, J. V. R. Ravindra and M. B. Srinivas "Minimizing Simulation Switching Noise (SSN) using Modified Odd/even Bus Invert Method" third IEEE international Workshop on Electronic design, test and applications (DELTA'06) 0-7695-2500-8/05\$20.00 2005 IEEE.
9. K. S. Sainarayanan C. Raghunandan, J.V. Ravindra, M. B. Srinivas "Efficient Spatial-Temporal Coding Scheme for Minimizing Delay in Interconnects" 2006 IEEE.
10. K. S. Sainarayanan, J. V. R. Ravindra, Kiran T. Nath, M. B. Srinivas "Coding for Minimizing Energy in VLSI Interconnects". 2006 IEEE.
11. Ketan N. Patel, Igor L. Markov "Error – Correction and Crosstalk Avoidance in DSM buses" IEEE transactions on VLSI systems, VOL.12, 12, NO.10, October 2004.
12. K.S. Sainarayan, J. V. R. Ravindra, M.B. Srinivas "A Novel, coupling driven, low power bus coding technique for minimizing capacitive crosstalk in VLSI interconnects" 2006 IEEE.
13. Lingamneni Avinash, M. Krithi Krishna, M. B. Srinivas "A novel encoding scheme for delay and energy minimization in VLSI Interconnect with Built-In Error Detection" IEEE Computer Society Annual Symposium on VLSI 2008.
14. Mircea R. Stan, Wayne P. Burleson "Bus- Invert Coding for Low-Power I/O" IEEE Transaction on VLSI 1995.
15. M. Ismail and N. tan "Modeling Techniques for Energy-Efficient System-on-chip signaling" IEEE Circuits and Magazine 2003.
16. R. P. Singh & S. D. Sapre "Communication Systems Analog & Digital" Chapter-11 coding page (505-519).



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 1 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Modification of Support Vector Machine for Microarray Data Analysis

By Vrushali Dipak Fangal & Dr. Sk. Sarif Hassan

Indian School of Mines, India

Abstract - The role of protuberant data analysis in selection of certain genes having distinctive level of activities between conditions of interest i.e diseased gene and normal genes is very significant. Now-a-days it is become a standard in gene analysis that microarray of DNA is a crucial data preparation step in systemization and other biological analysis. We consider the problem of constructing an accurate prediction rule for separating the different labels of genes in microarray gene expression data. Use of SVM in such data analysis is not new but it is not up to the mark we desire. So in this manuscript, we have tried to modify Support Vector Machine (SVM) for better accuracy in cancer genes systemization. Here we have modified SVM to account for gene redundancy and keep a check on it. In the other approach, instead of keeping bias a constant in SVM, we have tried to modify SVM by bias variation which we call as Orthogonal Vertical Permutator (OVP).

Keywords : support vector machine, microarray, redundancy, orthogonal vertical permutator.

GJCST-A Classification : C.1.2



Strictly as per the compliance and regulations of:



Modification of Support Vector Machine for Microarray Data Analysis

Vrushali Dipak Fangal^a & Dr. Sk. Sarif Hassan^c

Abstract - The role of protuberant data analysis in selection of certain genes having distinctive level of activities between conditions of interest i.e diseased gene and normal genes is very significant. Now-a-days it is become a standard in gene analysis that microarray of DNA is a crucial data preparation step in systemization and other biological analysis. We consider the problem of constructing an accurate prediction rule for separating the different labels of genes in microarray gene expression data. Use of SVM in such data analysis is not new but it is not up to the mark we desire. So in this manuscript, we have tried to modify Support Vector Machine (SVM) for better accuracy in cancer genes systemization. Here we have modified SVM to account for gene redundancy and keep a check on it. In the other approach, instead of keeping bias a constant in SVM, we have tried to modify SVM by bias variation which we call as Orthogonal Vertical Permutator (OVP).

Keywords : support vector machine, microarray, redundancy, orthogonal vertical permutator.

1. INTRODUCTION

The theory of support vector machines (SVMs), which is based on the idea of structural risk minimization (SRM), is a new classification technique and has drawn much attention on this topic in recent years (Burges, 1998; Cortes and Vapnik 1995; Vapnik, 1995, 1998). The good generalization ability of SVMs is achieved by finding a large margin between two classes (Bartlett and Shawe-Taylor, 1998; Shawe-Taylor and Bartlett, 1998). In many applications, the theory of SVMs has been shown to provide higher performance than traditional learning machines (Burges, 1998) and has been introduced as powerful tools for solving classification problems. Since the optimal hyper-plane obtained by the SVM depends on only a small part of the data points, it may become sensitive to noises or outliers in the training set (Boser et al., 1992; Zhang, 1999).

To solve this problem, one approach is to do some preprocessing on training data to remove noises or outliers, and then use the remaining set learn the decision function (Cao et al., 2003). This method is hard to implement if we do not have enough knowledge about noises or outliers. In many real world applications, we are given a set of training data without knowledge

about noises or outliers. There are some risks to remove the meaningful data points as noises or outliers.

Support Vector Machines have gained much attention in recent years due to their better predictability and ability to theoretically project any data to infinite dimension. It works on the simple basis of separating classes using a hyper plane.

In present scenario, any classification method has to deal with thousands of genes provided by micro array data. This is real test for and classification method. Neural networks have shown high potential in dealing with huge amount of data but it cannot overcome redundancy problem. Many highly correlated genes play similar role in classification while many of them could be omitted. Support Vector Machines also could not account for this problem in current form. In SVM, the weights of any two highly correlated features will be quite near and thus both can play significant role in classification. It is a major hindrance in feature selection.

In this paper, we present two different approaches for improvement of classification accuracy for linear SVMs. In the first method, redundancy control has been targeted for improve the classification rate. For checking a control on redundancy an matrix 'A' has been introduced in the optimization problem. This matrix keeps a check on weight of a feature according to its correlation with other features. It will be discussed in details in later section of paper.

The second method is also an approach to improve the classification performance of linear SVM. It is based on adjustment of bias value in SVM. The results have encouraged us for further probe.

In the paper, Section 2 describes the architecture of normal Support vector machine. Section 3 compares the architecture of normal SVM and modified SVM for controlling the redundancy. Section 4 describes the other method for improving the classification accuracy called Orthogonal Vertical Permutator. Section 5 and 6 discusses experimentation and the results obtained respectively.

II. SVM AND PROPOSED MODIFICATION IN SVM

a) SVM

Support Vector Machines (SVM's) are learning methods used for binary classification of data. The basic idea is to find a hyper-plane separating n-dimensional

Author a : Department of Computer Science and Engineering, Indian School of Mines, Dhanbad, Jharkhand-826004, India.

E-mail : vrush2elu@ismu.ac.in

Author c : Institute of Mathematics & Applications Bhubaneswar, India.

E-mail : sarimif@gmail.com

data perfectly into two classes. Since the example data is often not linearly separable, SVM introduce a "kernel induced feature space" which casts data into a higher dimensional space where data is separable. SVM plays a major role in eliminating computational complexity and over fitting (Crammer, Koby, 2001, Drucker, Harris, 1996, Ferris, Michael C, 2002 and T. S. Furey et al. 2000).

We are given l training samples $\{x_i, y_i\}$, $i = 1, \dots, l$, where each sample has d inputs ($x_i \in \mathbb{R}^d$), and a class label with one of two values ($y_i \in \{-1, 1\}$). Now, all hyper-planes in $\mathbb{R}^d$ are parameterized by a vector (w), and a constant (b), expressed in the equation $w \cdot x + b = 0$ where w is orthogonal to the hyper-plane.

Given such a hyper-plane (w, b) that separates the data, this gives the function

$$f(x) = \text{sign}(w \cdot x + b)$$

which correctly systemizes the training. However, a given hyper-plane represented by (w, b) is equally expressed by all pairs $\{\lambda w, \lambda b\}$ for $\lambda \in \mathbb{R}^+$. So we define the canonical hyper-plane to be that which separates the data from the hyper-plane by a distance of at least 1.

That is, we consider those that satisfy:

$$x_i \cdot w + b > +1 \text{ when } y_i = +1 \text{ and } x_i \cdot w + b < -1$$

when or more compactly: $y_i(x_i \cdot w + b) > 1 \forall i$.

We can frame this as an optimization problem as:

Minimize in (w, b): $\|w\|$ subject to (for any $i = 1, \dots, n$) $y_i(w \cdot x_i - b) \geq 1$

b) Modified SVM

Before we start the modification over the existing SVM let us understand the method of generating the matrix A .

i. Generation of 'A' Matrix

1. On basis of property of features like correlation or mutual information.
2. Using a function of importance or unique data points.
3. We may also use something like gradient descent method.

$$\text{If } w_i = \frac{w_i}{1 + \beta_i};$$

Then $E = (\text{Calculated Output} - \text{Actual Output})^2$;

$$\text{So, } E = (w^T x_i + b - O_i)^2$$

$$\text{Then, } \frac{\partial E}{\partial \beta} = \sum_{i=1}^l 2[w^T x_i + b - O_i] \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix}$$

$$a_d = \frac{x_{id} \cdot (-w_d)}{(1 + \beta_d)^2}$$

Consider the optimization problem in SVM. We introduce a matrix 'A' of order $n \times n$ where 'n' is number of features. If we minimize this optimization problem, the weight vector obtained is different from normal SVM

weight vector and we can keep a check on redundancy depending on 'A' matrix.

$$\text{min } \frac{1}{2} \|w\|^2 + C \sum_{i=1}^l \xi_i + \frac{1}{2} w^T A w$$

$$\text{s.t. } y_i(w^T x_i + b) \geq 1 - \xi_i \text{ where } A =$$

$$\begin{pmatrix} a_{11} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & a_{nn} \end{pmatrix} \text{ is a diagonal matrix.}$$

Introducing Lagrange's multiplier and converting to dual form

$$\phi(w, b, \xi, \alpha, \beta) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^l \xi_i + \frac{1}{2} w^T A w - \sum_{i=1}^l \alpha_i [y_i(w^T x_i + b) - 1 + \xi_i] + \sum_{i=1}^l \beta_i \xi_i$$

$$\frac{\partial \phi}{\partial b} = 0 \Rightarrow \sum_{i=1}^l \alpha_i y_i = 0$$

$$\frac{\partial \phi}{\partial \xi} = 0 \Rightarrow \alpha_i + \beta_i = C$$

$$\frac{\partial \phi}{\partial \xi} = 0 \Rightarrow w = B^{-1} \sum \alpha_i y_i x_i = \frac{\sum \alpha_i y_i x_i}{1 + \alpha_i}$$

Change in Hessian Matrix is-

$$H = y_i y_j x_i x_j$$

$$H = y_i y_j \frac{x_i}{1 + \alpha_i} x_j$$

ii. Comparing Architecture of SVM with Modified SVM

The layout of normal SVM has been shown below. A separating hyper-plane in canonical form must satisfy the following constraints,

$$y_i[\langle w, x_i \rangle + b] \geq 1, i = 1, \dots, l.$$

The distance $d(w, b; x)$ of a point x from the hyper-plane (w, b) is

$$d(w, b; x) = \frac{|\langle w, x^i \rangle + b|}{\|w\|}$$

$$p(w, b) = \min_{x^i, y^i = -1} d(w, b; x^i) + \min_{x^i, y^i = 1} d(w, b; x^i)$$

$$p(w, b) = \min_{x^i, y^i = -1} \frac{|\langle w, x^i \rangle + b|}{\|w\|} + \min_{x^i, y^i = 1} \frac{|\langle w, x^i \rangle + b|}{\|w\|}$$

$$p(w, b) = \frac{1}{\|w\|} \left(\min_{x^i, y^i = -1} |\langle w, x^i \rangle + b| + \right.$$

$$\left. \min_{x^i, y^i = 1} |\langle w, x^i \rangle + b| \right)$$

$$p(w, b) = \frac{2}{\|w\|}$$

Hence, the hyper-plane that optimally separates the data is the one that minimizes

$$\phi(w) = \frac{1}{2} \|w\|^2$$

This is solved by using Lagrange's multipliers.

$$\phi(w, b; \alpha) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^l \alpha_i (y_i [w \cdot x^i + b] - 1),$$

Where α are the Lagrange multipliers. The Lagrangian has to be minimised with respect to w , b and minimised with respect to $\alpha \geq 0$. Classical Lagrangian duality enables the primal problem,

$$\max_{\alpha} W(\alpha) = \max_{\alpha} (\min \phi(w, b; \alpha))$$

The minimum with respect to w and b of the Lagrangian, ϕ , is given by,

$$\frac{\partial \phi}{\partial b} = 0 \Rightarrow \sum_{i=1}^l \alpha_i y_i = 0$$

$$\frac{\partial \phi}{\partial w} = 0 \Rightarrow w = \sum_{i=1}^l \alpha_i y_i x_i$$

Hence, the dual problem is

$$\max_{\alpha} W(\alpha) = \max_{\alpha} - \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle$$

$$> + \sum_{k=1}^l \alpha_k$$

And hence the solution to the problem is given by

$$\alpha^* = \arg \min_{\alpha} \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j \langle x_i, x_j \rangle - \sum_{k=1}^l \alpha_k$$

With constraints,

$$\alpha_i \geq 0 \quad i = 1, \dots, l$$

$$\sum_{j=1}^l \alpha_j y_j = 0$$

This equation can be represented as a quadratic form.

iii. Orthogonal Vertical Permutator

Orthogonal Vertical Permutator is a reformation of SVM. In OVP, we vary the bias value of SVM which results in vertical permutations of the hyper-plane resulting from SVM. This section of paper focuses on the bias value 'b' in SVM framework. Its theoretical

inspiration is being discussed in following section. Bias is the constant term which is added in decision making equation.



Figure 1 : Depicting the concept of OVP in SVM

a. Adjustment of Bias Value in SVM

Consider two concentric circles with each circle representing same class as in Fig. 1. Here, we compare the correct output rate of the output generated by SVM and OVP.

Let the radius of inner circle be 'r' and radius of outer circle be n times 'r' i.e 'nr'. Consider each point on the circle represents a sample.

b. Correct Classification by SVM

$$\text{Correct classification} = \pi r + \pi n r$$

$$A = (n + 1)\pi r \text{ which is half circumference of each circle}$$

c. Correct Classification by OVP

Inner circle is classified correctly.

$$\text{Correct classification} = 2\pi r + \text{Correct classification of outer circle}$$

We need to find θ to find the correct classification of outer circle. $\theta = \text{Arc Radius}$

$$\theta = \frac{\text{Arc}}{\text{Radius}}$$

$$\frac{\pi}{2} - \frac{\theta}{2} = \sin^{-1}\left(\frac{r}{nr}\right)$$

$$\frac{\pi}{2} - \frac{\theta}{2} = \sin^{-1}\left(\frac{1}{n}\right)$$

Therefore,

$$\text{Arc} = \theta \times \text{Radius}$$

$$\text{Arc} = \left(\pi - 2\sin^{-1}\left(\frac{1}{n}\right)\right) \times nr$$

$$\text{Arc} = \left(\pi nr - 2nr\sin^{-1}\left(\frac{1}{n}\right)\right)$$

Thus, the total correct output of by OVP is-

$$B = 2\pi r + \pi nr - 2nr\sin^{-1}\left(\frac{1}{n}\right)$$

Difference in the correct output by SVM and OVP is

$$B - A = \left(2\pi r + n\pi r - 2nr \sin^{-1} \frac{1}{n} \right) - (n+1)\pi r$$

$$B - A = \left(\pi r - 2nr \sin^{-1} \frac{1}{n} \right) \geq 0$$

Since, SVM generates the hyper-plane with the best possible slope, here we have adjusted the bias value to shift this plane using minimization of classification error of both classes. Therefore it can be seen that classification error is less in the later case as compared to the normal SVM. For realizing this plane, an approach similar to gradient descent is used. Bias value is changed by a fraction of its current value depending on the minimization of error.

The error in this case is defined in a different way than in usual case.

Normally error is defined as,

$$\text{Error} = \frac{\text{Total Misclassification of Class(-1) and Class(1)}}{\text{Total Number of Samples}}$$

But in this case, we have defined Error as

$$\text{Error} = \frac{\text{Total Misclassification of Class(-1)}}{\text{Total Number of Samples of Class(-1)}} + \frac{\text{Total Misclassification of Class(+1)}}{\text{Total Number of Samples of Class(+1)}}$$

Both error rates are quite different. In first case the, each error has absolute importance and is equally important. But in the other case, the error rate for each class is different and its importance is related to the number of samples in its class. In second case, one class may be classified to very high accuracy at the expense of the other. It leads to higher probability of accuracy rate for one class.

SVM is used to categorize datasets into binary data. All the hyper-planes separating the data into two groups are orthogonal to vector w . The variation of bias gives rise to various permutations of the hyper-planes along the vertical. Our model gives rise to vertical permutations of the orthogonal hyper-planes and hence the Orthogonal Vertical Permutator is named.

III. RESULT AND ANALYSIS

This section gives the comparison of percentage accuracy of SVM with modified SVM in Table-I against the sigma values of A matrix. It can be inferred that the modification offers a better accuracy over SVM. The second table depicts a comparison of percentage of accuracy of SVM and OVP-SVM. The percentage accuracy with shifted bias value is better than normal bias value. Figure 2 and figure 3 gives the graphical representation of application of SVM and OVP-SVM on concentric circle dataset and spiral dataset respectively.

A Matrix	Modified SVM	SVM
All Sigmas	75.85	75.25
Sigma>0.75	75.05	74.45
Selected Sigmas	76.3	75.4

Table 1 : Comparing Accuracy Results of SVM and Modified SVM

Dataset	SVM	OVP-SVM
Concentric Circle	49.7%	63%
Spiral	49.8%	53.05%

Table 2 : Comparing Accuracy Results of SVM and OVP-SVM

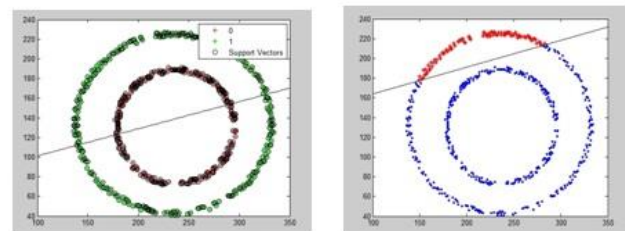


Figure 2 : Results of SVM and OVP-SVM on Concentric Circle Dataset

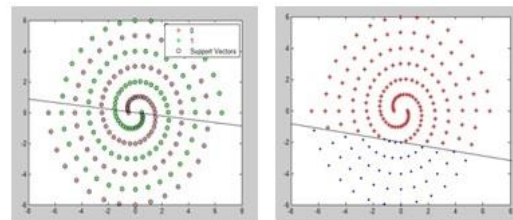


Figure 3 : Results of SVM and OVP-SVM on Spiral Dataset

The comparison analysis of both classification methods is done on the benchmark datasets. Each dataset is validated using double cross fold approach. Linear SVM is used for classification. Therefore only one parameter needs to be tuned i.e. the 'C' which accounts for soft margin classification. If training data was not provided separately then data was analyzed using 5 fold double cross validation. The data was divided into four parts of training data and one part of testing data. This training data was again five folded with four folds for actual training and one fold for parameter adjustment. All the datasets are available at UCI Machine Learning repository [6]. Table 1 shows the effect of change in matrix 'A' on sonar dataset. In rest of cases results are obtained by creating matrix 'A' is made by summing the correlation coefficient of a feature with rest of the features.

The experimentation of SVM with changed value of bias is performed on two datasets. The result of concentric circle dataset and of Spiral Dataset is shown in table 2.

The results obtained in both methods outperform the normal SVM and have different advantages. Modified SVM is immune to redundancy and OVP helps in improvising the classification accuracy of SVM and can be beneficial in multi class datasets. The processing was done on Matlab R2009.

REFERENCES RÉFÉRENCES REFERENCIAS

1. G. Eason, B. Noble, and I. N. Sneddon, "On certain integrals of Lipschitz-Hankel type involving products of Bessel functions," *Phil. Trans. Roy. Soc. London*, vol. A247, pp. 529–551, April 1955.
2. J. Clerk Maxwell, *A Treatise on Electricity and Magnetism*, 3rd ed., vol. 2. Oxford: Clarendon, 1892, pp. 68–73.
3. I. S. Jacobs and C. P. Bean, "Fine particles, thin films and exchange anisotropy," in *Magnetism*, vol. III, G. T. Rado and H. Suhl, Eds. New York: Academic, 1963, pp. 271–350.
4. Crammer, Koby; and Singer, Yoram (2001). "On the Algorithmic Implementation of Multiclass Kernel-based Vector Machines". *J. of Machine Learning Research* 2: 265–292. R. Nicole.
5. Drucker, Harris; Burges, Christopher J. C.; Kaufman, Linda; Smola, Alexander J.; and Vapnik, Vladimir N. (1997); "Support Vector Regression Machines", in *Advances in Neural Information Processing Systems 9*, NIPS 1996, 155–161, MIT Press.
6. Ferris, Michael C.; and Munson, Todd S. (2002). "Interior-point methods for massive support vector machines". *SIAM Journal on Optimization* 13 (3): 783–804.
7. T. S. Furey et al. (2000), support vector machine classification and validation of cancer tissue samples using microarray expressino data. *Bioinformatics*, 16 (10): 906-914.
8. Y. Yorozu, M. Hirano, K. Oka, and Y. Tagawa, "Electron spectroscopy studies on magneto-optical media and plastic substrate interface," *IEEE Transl. J. Magn. Japan*, vol. 2, pp. 740–741, August 1987 [Digests 9th Annual Conf. Magnetism Japan, p. 301, 1982].
9. M. Young, *The Technical Writer's Handbook*. Mill Valley, CA: University Science, 1989.



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 1 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Hardware Synthesis of Chip Enhancement Transformations in Hardware Description Language Environment

By Priyanka Saini, Adesh Kumar, Neha Singh & Dr. Anil Kumar Sharma
University of Petroleum and Energy Studies, India

Abstract - Human analyze different sight in daily life images to perceive their environment. More than 99% of the activity of human brain is involved in processing images from the visual cortex. A visual image is rich in information and can save thousand words. Many real world images are acquired with low contrast and unsuitable for human eyes to read, such as industrial and medical X-ray images. Image enhancement is a classical problem in image processing and computer vision. The image enhancement is widely used for image processing and as a preprocessing step in texture synthesis, speech recognition, and many other image/video processing applications. The main challenge is to transpose the validated algorithms into a language as hardware description languages. It is also the need that the input and output data files should be reshaped to match the binary content permitted into the hardware simulators. Research focuses on Simulation, Design and Synthesis of 2D and 3D Image enhancement chip in Hardware description language (HDL) Environment. The chip implementation of image enhancement algorithm is done using Discrete Wavelet Transformation (DWT) and Inverse Modified Discrete Cosine Transformation (IMDCT). Hardware chip modeling and simulation is done in Xilinx 14.2 ISE Simulator. Synthesis environment is carried out on Diligent Sparten-3E FPGA. . Image enhanced values are seen in the waveform editor of Modelsim software.

Keywords : VHDL- very high speed integrated circuit hardware description language, FPGA- field programmable gate array, HE – histogram equalization.

GJCST-A Classification : B.7.0



Strictly as per the compliance and regulations of:



RESEARCH | DIVERSITY | ETHICS

Hardware Synthesis of Chip Enhancement Transformations in Hardware Description Language Environment

Priyanka Saini ^α, Adesh Kumar ^σ, Neha Singh ^ρ & Dr. Anil Kumar Sharma ^ω

Abstract - Human analyze different sight in daily life images to perceive their environment. More than 99% of the activity of human brain is involved in processing images from the visual cortex. A visual image is rich in information and can save thousand words. Many real world images are acquired with low contrast and unsuitable for human eyes to read, such as industrial and medical X-ray images. Image enhancement is a classical problem in image processing and computer vision. The image enhancement is widely used for image processing and as a preprocessing step in texture synthesis, speech recognition, and many other image/video processing applications. The main challenge is to transpose the validated algorithms into a language as hardware description languages. It is also the need that the input and output data files should be reshaped to match the binary content permitted into the hardware simulators. Research focuses on Simulation, Design and Synthesis of 2D and 3D Image enhancement chip in Hardware description language (HDL) Environment. The chip implementation of image enhancement algorithm is done using Discrete Wavelet Transformation (DWT) and Inverse Modified Discrete Cosine Transformation (IMDCT). Hardware chip modeling and simulation is done in Xilinx 14.2 ISE Simulator. Synthesis environment is carried out on Diligent Spartan-3E FPGA. Image enhanced values are seen in the waveform editor of Modelsim software.

Keywords : VHDL- very high speed integrated circuit hardware description language, FPGA- field programmable gate array, HE – histogram equalization.

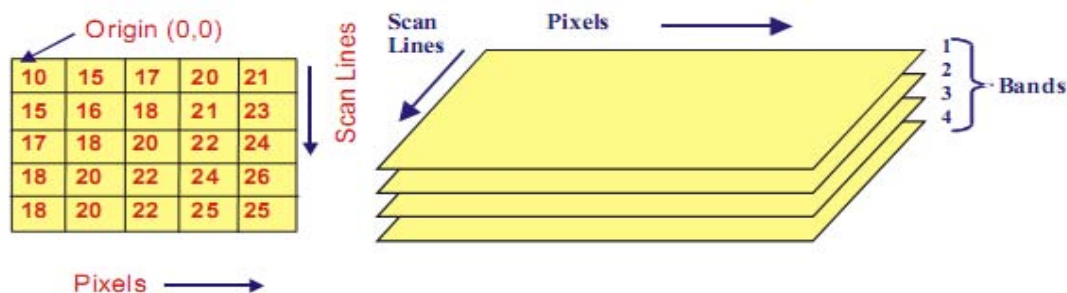


Figure 1 : Structure of a Digital Image and Multispectral Image [10]

Author ^α : M.Tech Scholar, Department of Electronics, Institute of Engineering and Technology, Alwar Rajasthan India.
E-mail : priyankasaini7@gmail.com

Author ^σ : Assistant Professor, Department of EEI, University of Petroleum and Energy Studies, Dehradun India.
E-mail : adeshmanav@gmail.com

Author ^ρ : Associate Professor, Department of Electronics, Institute of Engineering and Technology, Alwar India.
E-mail : nneha.singh01@gmail.com

Author ^ω : Professor, Department of Electronics, Institute of Engineering and Technology, Alwar India.
E-mail : aks_826@yahoo.co.in

I. INTRODUCTION

Pictures are the most common and convenient means of conveying [4] or transmitting information. A picture is worth a thousand words [3] [7]. Pictures concisely convey information about positions, sizes and inter-relationships between objects. Human recognize the images as object which are represented in spatial information [1] that we can recognize as objects. Human innate [1] [5] visual and mental abilities are good at deriving information from such images, because of 75% of the information received by human is in pictorial form. A digital remotely sensed image is typically composed of picture elements or pixels [2] [6] are located at the intersection of each row i and column j in each K bands of imagery. Associated with each pixel is a number known as Digital Number (DN) [2] or Brightness Value (BV) [3] that depicts the average radiance [19] of a relatively small area within a scene as shown in figure 1. A smaller number indicates low average radiance [7] from the area and the high number is an indicator of high radiant properties of the area. The size of this area effects the reproduction of details within the scene. As pixel size is reduced, more scene detail is presented in digital representation.

Image enhancement methods [1] [2] [8] are basically improving the perception or interpretability of information in images for human viewers. The reason of it is to provide better input for other automated digital image processing techniques [21]. The main objective of image enhancement is to change or modify the attributes of an image [14] can be suitable for different task and a specific observer can identity it to fulfill his requirements. During image enhancement process,

[21] one or more attributes of an image are modified. A specific task may be the choice of attributes and the ways they are modified are specific to a given task. Choices of image enhancement methods are subjected to observer-specific factors such as the human visual system [11]. The choice of image enhancement methods also depends on observer's experience and it will introduce a great deal of subjectivity into choice of image enhancement methods [12]. Image enhancement is used in the following cases:- enhancement of dark image [7], removal of noise and distortion from image [7] [12], enhancement of dark image and highlight the edges of the objects [14] in an image. Different Transformations can be applied to perform these operations.

Digital image processing can be implemented into digital chips. For example digital cameras [16] generally use dedicated digital image processing chips which are used to convert the raw data taken from image sensor into a colour image in a standard image file format. Further these images are used in digital cameras to improve their quality. A software program is used for the modification [16] in the image and can manipulate the images in different ways. Digital camera enable of viewing the histograms [12] [13] of images by which a photographer can understand rendered brightness range of each image shot more readily. Digital images play a very important role in our daily life applications such as satellite television, magnetic resonance imaging and computer tomography. An image is defined as an array [17], or a matrix, of square pixels [12] [17] arranged in rows and columns. These are also called picture elements. An image can be represented in 2D configuration for a 3 D scene [17]. An object can be represented by its numerical value by an image [15]. An image is said a 2D function that represents some characteristics such as intensity, colour and brightness [1] [16] of any scene. It can be defined as a two variable function $f(x,y)$ [17] projected in a plane where $f(x,y)$ defines the light intensity at particular point.

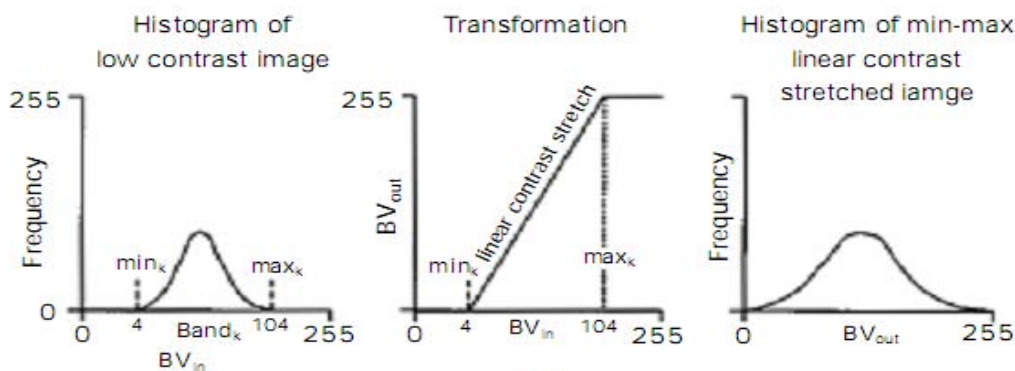
II. IMAGE ENHANCEMENT TRANSFORMATIONS

Image enhancement techniques [13] [14] improve the quality of an image as perceived by a human. These techniques are most useful because many satellite images [21] when examined on a color display give inadequate information for image interpretation. There is no conscious effort to improve the fidelity of the image with regard to some ideal form of the image. There exists a wide variety of techniques for improving image quality. The contrast stretch, edge enhancement, density slicing, and spatial filtering [3] are the more commonly used techniques. Image enhancement is attempted after the image is corrected for geometric and radiometric distortions [19]. Image enhancement methods are applied separately to each band of a multispectral image [17] [19]. Digital techniques [4] [12] have been found to be most satisfactory than the photographic technique for image enhancement, because of the precision and wide variety of digital processes. Image Enhancement Techniques are listed below:

- Contrast Enhancement Method
- Smoothing
- Brightening
- Intensity Transformation
- Discrete Wavelet Transformation
- IMDCT (Inverse Modified Discrete Cosine Transformation)

a) Contrast Enhancement Method

An Image is taken and its contrast is increased or decreased as per the enhancement requirements of the Image. The required contrast enhancement is achieved applying the Histogram Stretching [17] of the Image. There are two methods of image enhancement Linear and Nonlinear Contrast Stretch [15]. The grey values [4] in the original image and the modified image follow a linear relation in linear contrast method.



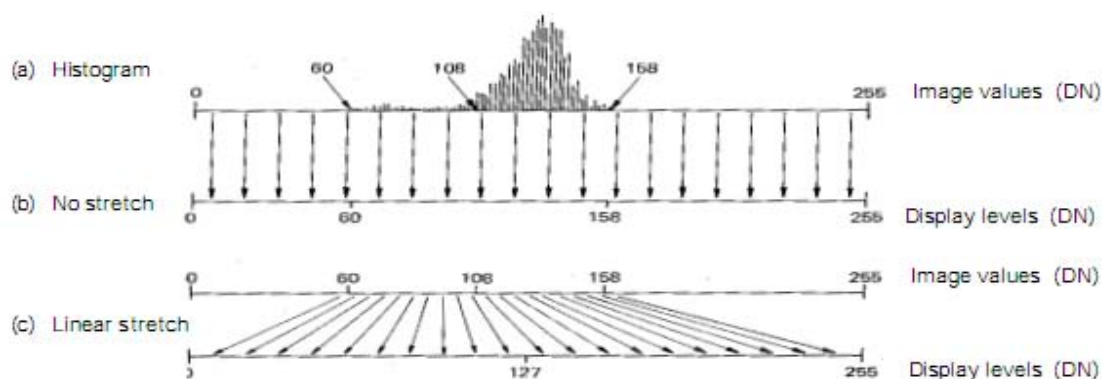


Figure 2: Linear Contrast Stretch [15]

A density number in the low range of the original histogram is assigned to extremely black and a value at the high end is assigned to extremely white. The remaining pixel values are distributed linearly between these extremes. The features or details that were obscure on the original image will be clear in the contrast stretched image [3]. Linear contrast stretch operation can be represented graphically as shown in Figure 2. To provide optimal contrast and color variation in color composites the small range of grey values in each band is stretched to the full brightness range [11] of the output or display unit. In Non-Linear contrast enhancement [17], the input and output data values follow a non-linear transformation [15]. The general form of the non-linear contrast enhancement is defined by $y = f(x)$, where x is the input data value and y is the output data value. The non-linear contrast enhancement techniques have been found to be useful for enhancing the color contrast between nearly classes and subclasses of a main class. A type of non linear contrast stretch involves scaling the input data logarithmically.

b) Smoothing

A noisy Image is taken and the noise removal [3] is done by applying a smoothing technique. The noise removal is achieved by using a mask which enables neighborhood pixel processing [15]. The aim of image smoothing is to diminish the effects of camera noise, spurious pixel values, [14] missing pixel values etc. There are many different techniques for image smoothing; we will consider neighborhood averaging and edge-preserving smoothing [15]. Each point in the smoothed image, $\hat{F}(x,y)$ is obtained from the average pixel value in a neighborhood of (x,y) in the input image.

For example, if we use a 3 x 3 neighborhood around each pixel we would use the mask

$$\begin{matrix} 1/16 & 1/16 & 1/16 \\ 1/16 & 1/16 & 1/16 \\ 1/16 & 1/16 & 1/16 \end{matrix}$$

Each pixel value is multiplied by 1/16, summed, and then the result placed in the output image. This mask is successively moved across the image until every pixel has been covered.

Neighborhood averaging or Gaussian smoothing will tend to blur edges because the high frequencies in the image are attenuated. An alternative approach is to use median filtering [3]. Here we set the grey level to be the median of the pixel values in the neighborhood of that pixel. The median m of a set of values is such that half the values in the set are less than m and half are greater. For example, suppose the pixel values in a 3x3 neighborhood are (10, 20, 20, 15, 20, 20, 20, 25, 100). If we sort the values we get (10, 15, 20, 20, 20, 20, 25, 100) and the median here is 20. Figure 3 (a) & (b) shows an image before smoothing and after smoothing.



Figure 3: (a) Image before smoothing & (b) After smoothing

c) Brightening

Enhancement techniques expand the range of brightness [15] values in an image so that the image can be efficiently displayed in a manner desired by the analyst. The density values in a scene are literally pulled farther apart, that is, expanded over a greater range. The effect is to increase the visual contrast [14] between two areas of different uniform densities. This enables the analyst to discriminate [21] easily between areas initially

having a small difference in density. Brightened Image [13] developed from original image by increasing every

pixel with a constant. Figure 4 (a) & (b) shows an image before brightening and after brightening.

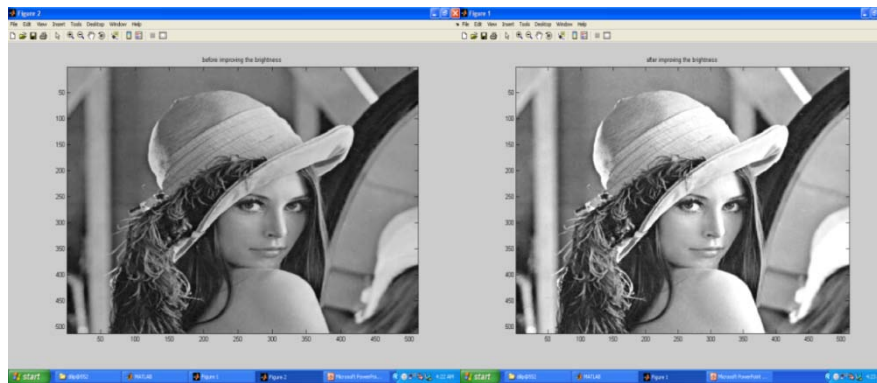


Figure 4 : (a) Image before bright and (b) After bright [15]

d) Intensity Transformation

The simplest form of the transformation T is when the neighborhood is $1 \times 1 \Rightarrow$ intensity transformation [13] [21]. Because they depend only on intensity values and not explicitly on the location of the pixel explicitly on the location of the pixel intensity, intensity transformation functions frequently are written as $s = T(r)$ where r denotes the intensity of f and s the intensity of g both at point (x, y) . for example, if $T(r)$ has the form in figure 5 (a), the effect of applying the transformation to every pixel of generate the corresponding pixels in g would be to produce an image

of higher contrast than the original by darkening the intensity levels below k and brightening the level above k . In this technique, sometimes called contrast stretching, values of r lower than k are compressed by the transformation function into a narrow range of s , toward black. The opposite is true for values of r higher than k . Otherwise how an intensity value r_0 is mapped to obtain the corresponding value s_0 . In the limiting case shown in figure 5 (b), $T(r)$ produces a two level image. A mapping of this type is called a Thresholding function [13]. By increasing the pixel size of any image we can enhance the image in x, y, z all the three directions.

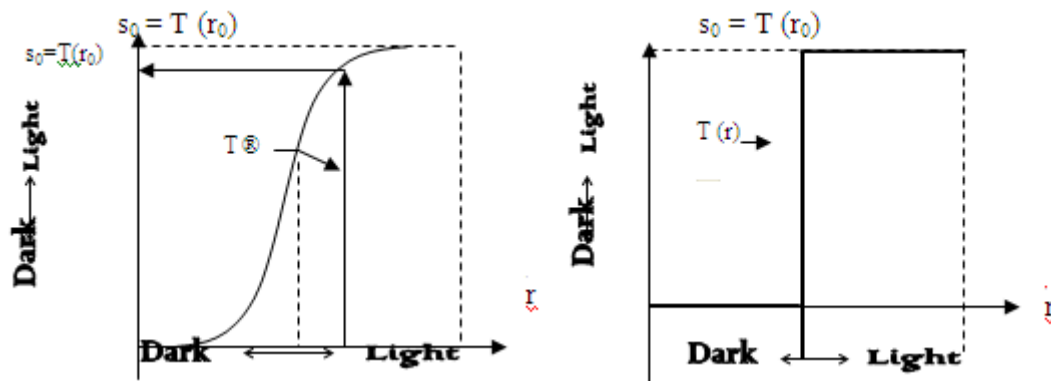


Figure 5 : Intensity Transformations (a). Contrast Stretching (b). Thresholding Function

e) Wavelet Transformation

The Discrete Wavelet Transform (DWT) [3] is a widely applicable image processing algorithm that is used in various applications. DWT decomposes an image by using scaled and shifted versions of a compact supported basis function called the mother wavelet, and provides a multi-resolution [18] representation of the image. Performing the DWT, modifying the transform coefficients, and performing the inverse transform (IDWT) [3] [19] of the modified coefficients is a promising method in signal and image

processing. It is based on the histogram equalization technique [19] to analyze the DWT and IDWT results.

f) IMDCT (Inverse Modified discrete cosine transform)

This transform is used for image compression and image enhancement [5]. It accepts 18 discrete values at one time. The 18-point IMDCT (block size 36) implementation is given by the following equation.

$$\widehat{x}_m = \frac{2}{N} \sum_{k=0}^{\left(\frac{N}{2}\right)-1} X_k \cdot \cos \left[\frac{\pi}{2N} (2k+1) \left(2m+1 + \frac{N}{2} \right) \right], \quad \text{with } m = 0, 1, 2, \dots, N-1$$

Generally, it is a lapped transform, the recovered data sequence $\{\widehat{x}_m\}$ does not correspond to the original data sequence $\{x_m\}$. To obtain the correct $\{x_m\}$ the outputs of consecutive transforms have to be combined. It can be seen that $N/2$ (non redundant) input values result in N output values (of course the MDCT [5] reads N input values and results in $N/2$ output values). Since it is not completely clear whether Equation 1

should be called an N -point IMDCT [5] or an $N/2$ -point IMDCT, in the following we shall identify these transforms given the number of inputs. Considering an 18-point IMDCT that delivers 36 output values, thus length N will be 36. Considering a case for $N=36$, we start from an 18 values input sequence: $\{X_0, X_1, \dots, X_{17}\}$. The output of rotational block is given by

$$a_n = X_n \cos \left[\frac{\pi}{2N} (2n+1) \right] + X_{N/2-1-n} \sin \left[\frac{\pi}{2N} (2n+1) \right]$$

$$b_n = X_n \sin \left[\frac{\pi}{2N} (2n+1) \right] - X_{N/2-1-n} \cos \left[\frac{\pi}{2N} (2n+1) \right]$$

$$\text{Here } n = 0, 1, 2, \dots, \frac{N}{4} - 1$$

The left most 'combine and shuffle'-block is thus nothing more than a reverse ordering of the second half of the input data.

III. DESIGN METHODOLOGY & FLOW

Figure 6 shows a flow chart over the design process when a design is implemented into an FPGA [13]. This flow was followed with all designs in this project and so became an important structure in the project plan. For those that is not familiar with these concepts a short description will follow. For more case specific see all the steps listed below at front end design. The Spartan 3E [23] [24] starter kit provides easy way to test the various programs in the FPGA itself, by dumping the 'bit' file of the designed program in Xilinx software into the FPGA and then observing the output. The Spartan 3E FPGA board comes built in with many peripherals that help in the proper working of the board and also in interfacing the various signals to the board itself. Some of the peripherals included in the Spartan 3E FPGA board include: 2-line, 16-character LCD screen used for display the output, PS/2 mouse [23] or keyboard port can be connected to the FPGA, VGA display port [24] used to display various encoded images via a screen.

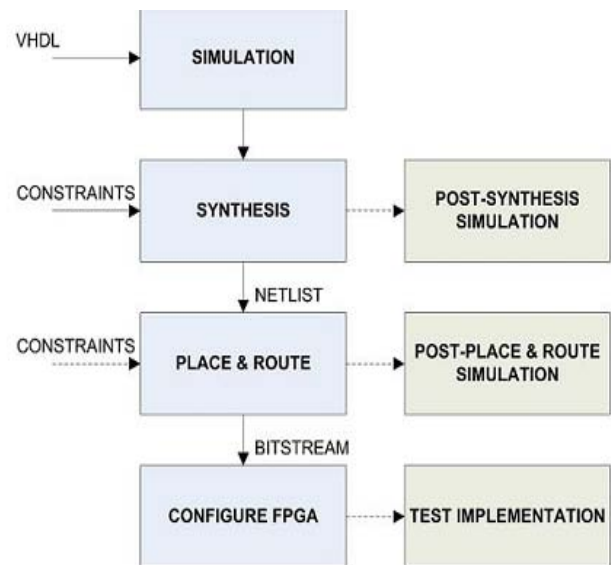


Figure 6 : FPGA Design Project Flow [22]

The image encoding would be done by the FPGA via the aid of the program and then the encoded image would be displayed on the screen. Two 9-pin RS-232 [23] ports help in the transmission of serial data to and from the FPGA board, 50 MHz clock oscillator is the system clock which helps in giving the clock signal to the various events taking place within the FPGA and the various programs that require clock for their working, A Digital clock manager [23] [24] can also be used to reduce the frequency of the system clock so that it is useful for various other purposes which need smaller clock frequency. On-board USB-based FPGA [24] download and debug interface is also in the Spartan-3E

kit in which the programmable file is dumped into the FPGA via the USB based download cable. Hence it is very much helpful in the testing of the programs whether they are working correctly or not, eight discrete LEDs can be interfaced to glow when a particular output becomes high. Hence the LEDs can be interfaced to show the output of a single bit. Four slide switches and four push-button switches are used to give the inputs to the FPGA board. They can also act as the reset switches for the various programs. Kit also has four-output, SPI-based on board Digital-to-Analog Converter (DAC) on board which is to be interfaced to give the analog output to the digital data values.

IV. SIMULATION RESULTS

Figure 7 (a) and (b) shows the snapshot results for image enhancement algorithm using DWT for 2D and 3D images respectively. Similarly, Figure 8 (a) and (b) shows the snapshot results for image enhancement algorithm using IMDCT for 2D and 3D images respectively. Simulation result is shown, considering 9 x 9 image sizes for 2D and 9 x 9 x 9 image for 3D.

Step Input 1 : reset =1, clk is applied for synchronization and then run.

Step Input 2 : reset =0, clk is applied for synchronization.

In the modelsim waveforms *image_in_x_axis*, *image_in_y_axis*, *image_in_z_axis* represents the integer value of image at discrete points at X, Y and Z axis respectively which is a matrix of 9 pixels vales. *Sample_x_axis*, *Sample_y_axis*, *image_z_axis* represents the integer value of image in discrete samples at X, Y and Z axis respectively by which image should be enhanced. It is also a matrix of 9 pixels vales corresponding to each input value of image. *image_out_x_axis*, *image_out_y_axis*, *image_out_z_axis* represents the corresponding results of image enhancement at X, Y and Z axis respectively within a matrix of 9 pixels vales. *P-state* and *n_state* are the present state and next state to develop the chip using Finite State Machine (FSM).

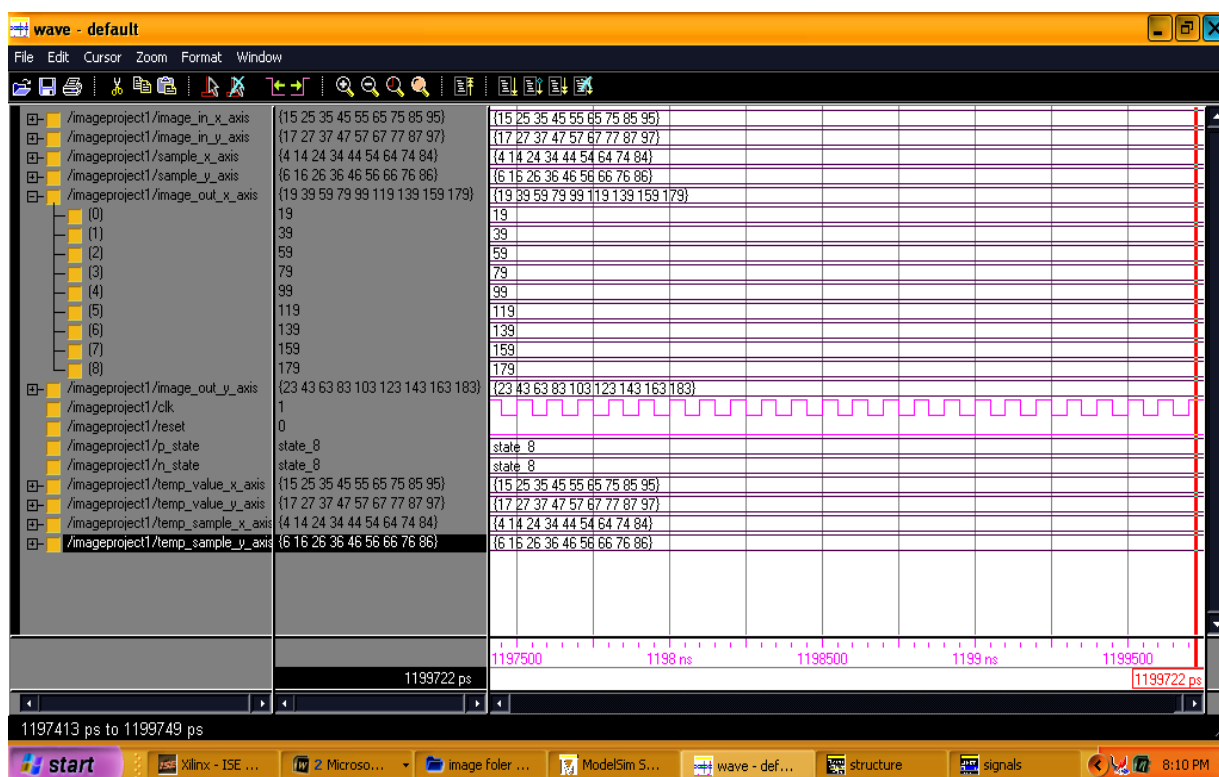


Figure 7 (a) : Modelsim waveform for 2D image enhancement chip using DWT

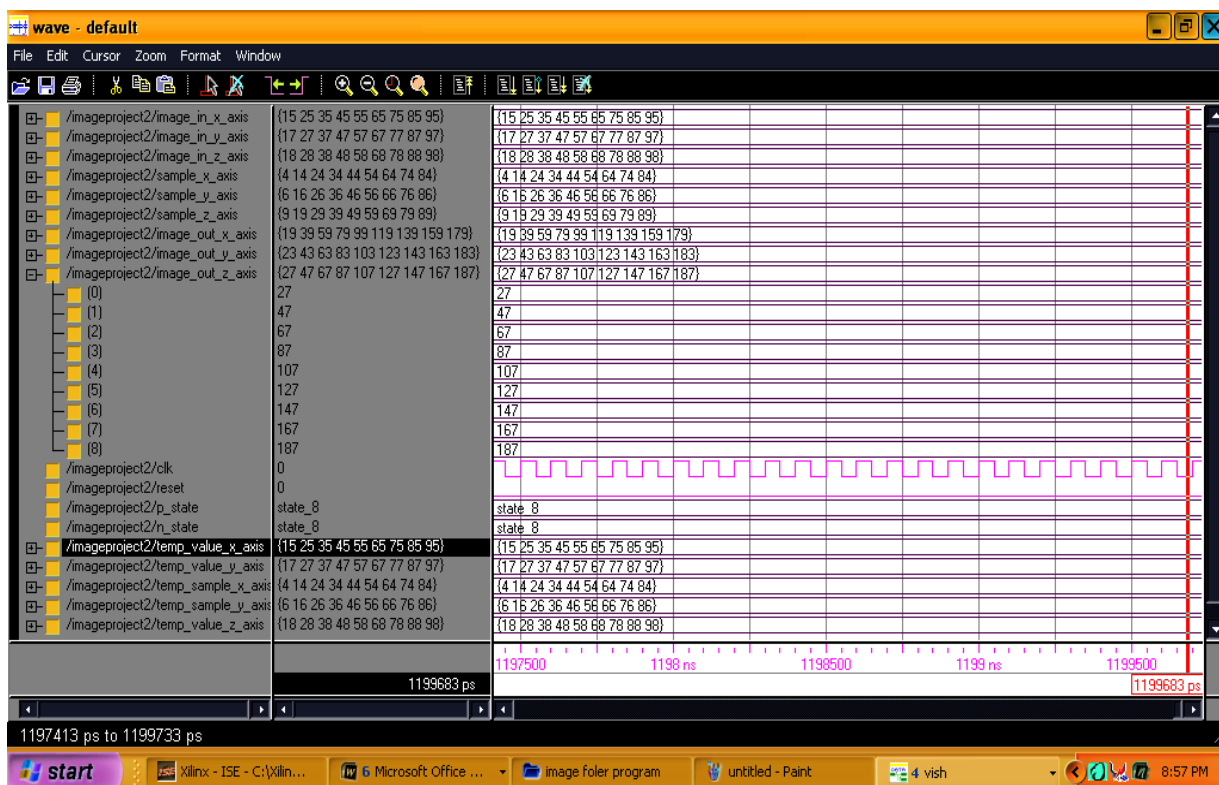


Figure 7 (b) : Modelsim waveform for 2D image enhancement chip using DWT

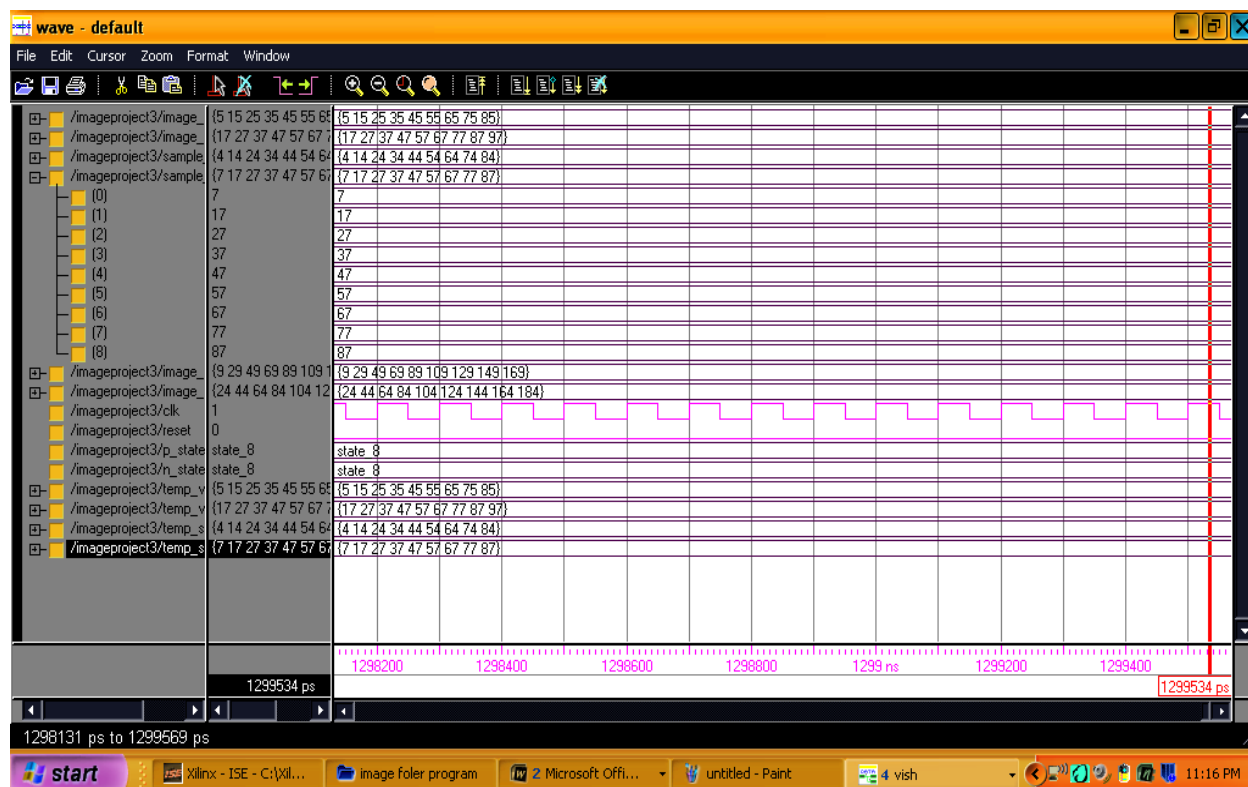


Figure 8 (a) : Modelsim waveform for 2D image enhancement chip using IMDCT

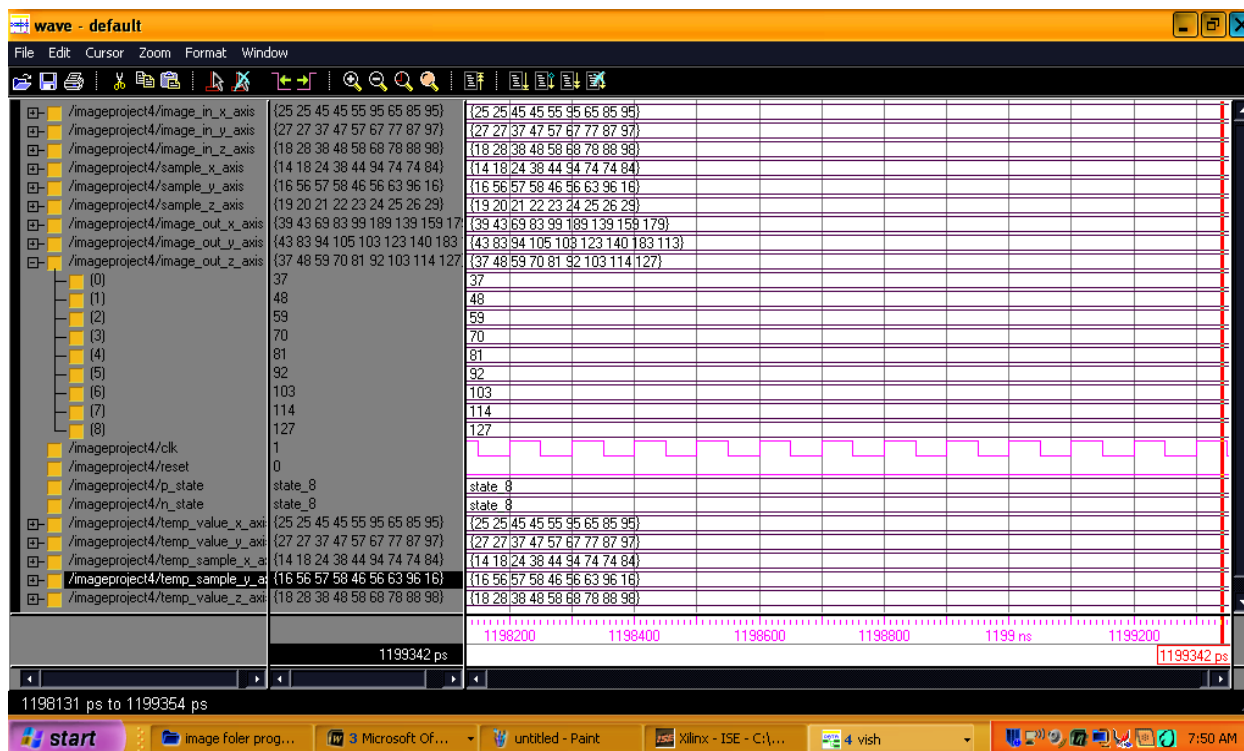


Figure 8 (b) : Modelsim waveform for 3D image enhancement chip using IMDCT

V. DEVICE UTILIZATION AND TIMING SUMMARY

Device utilization report is the report of used device hardware in the implementation of the chip and

timing report is the minimum and maximum time to reach the output. Table 1 and 2 shows the detail of device utilization for 2D and 3D images using DWT and IMDCT respectively.

Selected Device: xc3s250e-5pq208

Table 1 : Device utilization for 2D and 3D image using DWT

Device Part	2D Image (DWT)			3D Image (DWT)		
	Used	Available	Utilization	Used	Available	Utilization
Number of Slices	1294	2448	52 %	1934	2448	79 %
Number of Flip Flops	20	4896	0 %	20	4896	0 %
Number of 4 input LUTs	1164	4896	23 %	1740	4896	35%
Number of bonded IOBs	149	158	94 %	149	154	97 %
Number of GCLKs	8	24	33 %	8	24	33 %

Table 2 : Device utilization for 2D and 3D image using IMDCT

Device Part	2D Image (IMDCT)			3D Image (IMDCT)		
	Used	Available	Utilization	Used	Available	Utilization
Number of Slices	1031	2448	42 %	1541	2448	62 %
Number of Flip Flops	20	4896	0 %	20	4896	0 %
Number of 4 input LUTs	705	4896	14 %	1055	4896	21%
Number of bonded IOBs	149	158	94 %	149	154	97 %
Number of GCLKs	8	24	33 %	8	24	33 %

Timing Summary

Speed Grade:- 5

Table 3 : Timing details for 2D and 3D image using DWT and IMDCT

Parameters	Utilization			
	2D Image		3D Image	
	Using DWT	Using IMDCT	Using DWT	Using IMDCT
Minimum Period	2.2012 ns	2.0071 ns	2.2511 ns	2.0978 ns
Maximum Frequency	715 MHz	715 MHz	715 MHz	715 MHz
Minimum input arrival time before clock	6.115ns	2.055ns	6.363ns	2.157ns
Maximum output required time after clock	4.179ns	4.179ns	4.179ns	4.179ns
Memory Utilization	141048 kB	142072 kB	155384 kB	156408 kB

Device utilization summary shows that number of slice utilization is reduced to 10 % in hardware chip design for 2D image using IMDCT, for 3D there is 16 % reduction in number of slices using IMDCT. There is a reduction of 9 % and 14 % in number of LUTs for 2D and 3D image chip using IMDCT respectively. Numbers of Flip Flops, bounded I/O, Number of GCLKs are same for both. Memory utilization shows an increment of 0.72 % for 2D and 0.65 % for 2D using IMDCT. Timing summary shows, there is 9 % reduction in minimum period for 2D and 7 % for 3D using IMDCT. There is great reeducation in Minimum input arrival time before clock which is 66 % change for 2D and 63 % for 3D using IMDCT method. The value of Maximum Frequency and Maximum output required time after clock are same.

VI. CONCLUSION

Image enhancement chip develop for 2D and 3D image is done using DWT and IMDCT transformations. Hardware parameter shows that there is 10 % reduction in number of slices in chip design for 2D image and 16 % reduction for 3D using IMDCT. There is great reeducation in Minimum input arrival time before clock which is 66 % change for 2D and 63 % for 3D using IMDCT method. Such applications are used in digital camera, satellite imaging, digital watermarking, X-rays, medical imaging, facial reorganization, Optical character reorganization and authenticity. This work is a significant effort towards total digitization of image processing and would surely prove a boon for VLSI design industry.

REFERENCES RÉFÉRENCES REFERENCIAS

1. T. Moreo, P. N. Lorente, F. S. Valles, J. S. Muro, C. F. Andrés, Experiences on developing computer vision hardware algorithms using Xilinx system generators, *Microprocessors and Microsystems* 29, (2005).
2. A. Sumathi, and Dr. R. S. D. Wahida Banu "Digital Filter Design using Evolvable Hardware Chip for Image Enhancement" *IJCSNS International Journal of Computer Science and Network Security*, VOL.6 No.5A, May 2006, page (201-210).
3. D. Dhanasekaran and K. Boopathy Bagan "HIGH SPEED PIPELINED ARCHITECTURE FOR ADAPTIVE MEDIAN FILTER," *European Journal of Scientific Research* ISSN 1450-216X Vol.29 No.4, pp. 454-460, 2004.
4. Gerasimos Louverdis, Ioannis Andreadis and Antonios Gasteratosn, "A NEW CONTENT BASED MEDIAN FILTER," 12th European Signal Processing Conference, pp. 1337-1340, September 2004.
5. Huibert J. Lincklaen Arriëns "Implementation of an 18-point IMDCT on FPGA, *HJLA, June-2005*", page (1-9).
6. Luliana CHIUCHISAN, Marius CERLINCA, Alin-Dan POTORAC, Adrian GRAUR "Image Enhancement Methods Approach using Verilog Hardware Description Language" 11th International Conference on DEVELOPMENT AND APPLICATION SYSTEMS, Suceava, Romania, May, 2012, page (144-148).
7. Juha Lemmetti, Juha Latvala, Hakan Öktem, Karen Egiazarian, and Jarkko Niittylahti "IMPLEMENTING WAVELET TRANSFORMS FOR X-RAY IMAGE ENHANCEMENT USING GENERAL PURPOSE PROCESSORS" Digital and Computer Systems Laboratory1, Signal Processing Laboratory2, Tampere University of Technology, Hermiankatu 12 C, 33720 Tampere, FINLAND, Page (1-4).
8. Karan Kumar, Aditya Jain, and Atul Kumar Srivastava "FPGA IMPLEMENTATION OF IMAGE ENHANCEMENT TECHNIQUES," *Proc. of SPIE* Vol. 7502 750208-7, September 2011.
9. Miguel A. Vega-Rodríguez, Juan M. Sánchez-Pérez and Juan A. Gómez-Pulido, "AN FPGA-BASED IMPLEMENTATION FOR MEDIAN FILTER MEETING THE REAL-TIME REQUIREMENTS OF AUTOMATED VISUAL INSPECTION SYSTEMS," 10th Mediterranean Conference on Control and Automation - MED2002, July 2002.
10. M. Chandrashekhar, U. Naresh Kumar, K. Sudershan Reddy "FPGA Implementation of High Speed Infrared Image Enhancement" *International Journal of Electronic Engineering Research*, Volume 1 Number 3 (2009) pp. 279-285.
11. N. R. Mokhtar, Nor Hazlyna Harun, M. Y. Mashor, H. Roseline, Nazahah Mustafa, R. Adollah, H. Adilah,

- N. F. Mohd Nasir "Image Enhancement Techniques Using Local, Global, Bright, Dark and Partial Contrast Stretching For Acute Leukemia Images" Proceedings of the World Congress on Engineering 2009 Vol. I WCE 2009, July 1-3, 2009, London, U.K.
12. Nitin Sachdeva and Tarun Sachdeva, "AN FPGA BASED REAL-TIME HISTOGRAM EQUALIZATION CIRCUIT FOR IMAGE ENHANCEMENT," IJECT Vol. 1, Issue 1, December 2010.
13. N. Shirazi, J. Ballagh, Put Hardware in the Loop with Xilinx System Generator for DSP, Xcell Journal, Fall 2003 pp 01-03.
14. Priyanka Saini, Adesh Kumar, Neha Singh " FPGA Implementation of 2D and 3D Image Enhancement Chip in HDL Environment" *International Journal of Computer Applications (0975 – 8887) Volume 62–No. 21, January 2013 page (24-31).*
15. Rafael C. Gonzalez and Richard E. Woods. Digital Image Processing. Addison-Wesley Publishing Company, 1992, chapter 4 page (212-235).
16. S. Sowmya, Roy Paily, "FPGA IMPLEMENTATION OF IMAGE ENHANCEMENT ALGORITHMS," International Conference Communications and Signal Processing (ICCSP), pp. 584 – 588, Feb. 2011.
17. S. Jayaraman, S Essakirajan, T VeeraKumar Book "Digital Image processing, Chapter-5 image Enhancement" Tata Mc Graw Hill Education Private Limited, New Delhi 2009, page (243-297).
18. Yadong Wu, Zhiqin Liu, Yongguo Han, Hongying Zhang, An Image illumination Correction Algorithm based on Tone Mapping, IEEE 3rd International Congress on Image and Signal Processing (CISP2010), pp 648-648.
19. Tarek M. Bittibssi, Gouda I. Salama, Yehia Z. Mehaseb and Adel E. Henawy "Image Enhancement Algorithms using FPGA" International Journal of Computer Science & Communication Networks, Vol. 2 issue 4 page (536-542).
20. T. Moreo, P. N. Lorente, F. S. Valles, J. S. Muro, C. F. Andrés, Experiences on developing computer vision hardware algorithms using Xilinx system generators, Elsevier Journal of Microprocessors and Microsystems 29, 2005 pp 411- 419.
21. Wang Bing-jian, Liu Shang-qian, Li Qing and Zhou Hui-xin, "A REAL-TIME CONTRAST ENHANCEMENT ALGORITHM FOR INFRARED IMAGES BASED ON PLATEAU HISTOGRAM," Infrared Physics & Technology, Elsevier, pp. 77-82, 2006.
22. Zhang, M.Z., Ngo, H.T., Asari, V.K.: Design of an Efficient Multiplier-Less Architecture for Multidimensional Convolution. Lecture Notes in Computer Science, Vol. 3740. Springer-Verlag, Berlin Heidelberg (2005) 65–78.
23. http://www.xilinx.com/support/documentation/data_sheets/ds312.pdf
24. http://www.xilinx.com/support/documentation/data_sheets/ds099.pdf
25. http://mosfet.isu.edu/classes/ee_cs_499_verilog_hdl/userguides/modelsimuser.pdf



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 1 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Parallel String Matching with Multi Core Processors-A Comparative Study for Gene Sequences

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

Andhra University, India

Abstract - The increase in huge amount of data is seen clearly in present days because of requirement for storing more information. To extract certain data from this large database is a very difficult task, including text processing, information retrieval, text mining, pattern recognition and DNA sequencing. So we need concurrent events and high performance computing models for extracting the data. This will create a challenge to the researchers. One of the solutions is parallel algorithms for string matching on computing models. In this we implemented parallel string matching with JAVA Multi threading with multi core processing, and performed a comparative study on Knuth Morris Pratt, Boyer Moore and Brute force string matching algorithms. For testing our system we take a gene sequence which consists of lacks of records. From the test results it is shown that the multicore processing is better compared to lower versions. Finally this proposed parallel string matching with multicore processing is better compared to other sequential approaches.

Keywords : *string matching; parallel string mathing; computing model, DNA, multicore processing.*

GJCST-A Classification : *B.7.1*



Strictly as per the compliance and regulations of:



Parallel String Matching with Multi Core Processors-A Comparative Study for Gene Sequences

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

Abstract - The increase in huge amount of data is seen clearly in present days because of requirement for storing more information. To extract certain data from this large database is a very difficult task, including text processing, information retrieval, text mining, pattern recognition and DNA sequencing. So we need concurrent events and high performance computing models for extracting the data. This will create a challenge to the researchers. One of the solutions is parallel algorithms for string matching on computing models. In this we implemented parallel string matching with JAVA Multi threading with multi core processing, and performed a comparative study on Knuth Morris Pratt, Boyer Moore and Brute force string matching algorithms. For testing our system we take a gene sequence which consists of lacks of records. From the test results it is shown that the multicore processing is better compared to lower versions. Finally this proposed parallel string matching with multicore processing is better compared to other sequential approaches.

Keywords : *string matching; parallel string mathing; computing model, DNA, multicore processing.*

I. INTRODUCTION

The crisis of finding exact or non-exact occurrences of a pattern P in a text T over some alphabet is a central difficulty of combinatorial string matching and has a variety of applications in many areas of computer science [1-3]. String searching algorithms can be accomplished in two ways:

1. Exact match, meaning that the passages returned will contain an exact match of the key input.
2. Approximate match, meaning that the passage will contain some part of the key word input [4-6].

Although the dramatic development of processor technology and other advances have reduced search response to negligible times, string matching problem still remains a useful area of research and development for a number of reasons. Initially, as the size of data continues to grow, sequence searches will become increasingly taxing on search engines. Secondly, the pattern matching still remains an integral part of faster matching algorithms, typically comprising

the final part of a search. Finally, researchers have to understand the classical methods of pattern matching to develop new efficient algorithms [7-10].

With the developments of new string matching techniques, efficiency and speed are the main factors in deciding among different options available for each application area. Each application area has certain special features that can be used by string matching technique best suited for that area [11-13]. This study implements a multithreading text searching approach to improve text searching performance at a multicore processing. The idea is to have more than one searcher thread that search the text from different positions. Since the required pattern may occur at any position, having multiple searchers is better than searching the text sequentially from the first character to the last one.

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 multicore computing models. This work suggests the most efficient algorithmic 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. RELATED WORK

Now a day's information retrieval attracted by the many researchers because of their importance in IT industry. So, this many researchers worked on this area since several years. In this paper we propose some of the techniques comparisons with multicore processing. In this section we discuss some previous techniques proposed by several authors' later section we will discuss about our actual procedure.

S.V. Raju et.al [14] proposes about grid computing in parallel string matching. Grid computing provides solutions for various complex problems. The function of the grid is to parallelize the string matching problem using grid MPI parallel programming method or loosely coupled parallel services on Grid. Parallel applications fall under three categories namely Perfect parallelism, Data parallelism and Functional parallelism, use data parallelism, and it is also called Single

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

E-mail : chinta.someswarao@gmail.com

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

Author ρ : Professor & HOD, Department of CSE, Kondagattu, Jagithyal, A.P., India.

Program Multiple Data (SPMD) method, where given data is divided into several parts and working on the part simultaneously.

Perfect Parallelism

Also known as embarrassingly parallel. An application can be divided into sets of processes that require little or no communication.

Data Parallelism

The same operation is performed on many data elements simultaneously. An example would be using multiple processes to search different parts of a database for one specific query.

Functional parallelism: Often called control parallelism. Multiple operations are performed simultaneously, with each operation addressing a particular part of the problem.

Result

Here it shows the performance of string matching algorithms namely execution-time and speedup improved.

HyunJin Kim and Sungho Kang [15] propose an algorithm that partitions a set of target patterns into multiple subgroups for homogeneous string matchers. Using a pattern grouping metric, the proposed pattern partitioning makes the average length of the mapped target patterns onto a string matcher approximately equal to the average length of total target patterns. The target architecture is based on a memory-based string matching with homogeneous string matchers. In a string matcher, N homogeneous finite state machine (FSM) tiles are contained. An FSM tile contains a maximum of s states and takes n bits of one character at each cycle. Target patterns are distributed and mapped onto C string matchers. Each state has $2n$ pointers for the next state based on an n -bit input.

Result

By adopting the pattern grouping metric, the proposed pattern group partitioning decreases the number of adopted string matchers by balancing the numbers of mapped target patterns between string matchers.

Daniel Luchaup, et.al [16] they propose a method to search for arbitrary regular expressions by scanning multiple bytes in parallel using speculation. They break the packet in several chunks, opportunistically scan them in parallel, and if the speculation is wrong, correct it later. They present algorithms that apply speculation in single-threaded software running on commodity processors as well as algorithms for parallel hardware.

Result

It is a speculative pattern matching method which is a powerful technique for low latency regular-expression matching. The method is based on three

important observations. The first key insight is that the serial nature of the memory accesses is the main latency-bottleneck for a traditional DFA matching. The second observation is that a speculation that does not have to be right from the start can break this serialization. The third insight, which makes such a speculation possible, is that the DFA-based scanning for the intrusion detection domain spends most of the time in a few hot states.

Hyun Jin Kim, et.al [17] propose a memory-efficient parallel string matching scheme. In order to reduce the number of state transitions, the finite state machine tiles in a string matcher adopt bit-level input symbols. Long target patterns are divided into sub patterns with a fixed length; deterministic finite automata are built with the sub patterns. Using the pattern dividing, the variety of target pattern lengths can be mitigated, so that memory usage in homogeneous string matchers can be efficient.

Result

The proposed DFA-based parallel string matching scheme minimizes total memory requirements. The problem of various pattern lengths can be mitigated by dividing long target patterns into sub patterns with a fixed length. The memory-efficient bit-split FSM architectures can reduce the total memory requirements. Considering the reduced memory requirements for the real rule sets, it is concluded that the proposed string matching scheme is useful for reducing total memory requirements of parallel string matching engines.

Charalampos S, et.al[18] they proposes that Graphics Processing Units (GPUs) have evolved over the past few years from dedicated graphics rendering devices to powerful parallel processors, outperforming traditional Central Processing Units (CPUs) in many areas of scientific computing. The use of GPUs as processing elements was very limited until recently, when the concept of General-Purpose Computing on Graphics Processing Units (GPGPU) was introduced. GPGPU made possible to exploit the processing power and the memory bandwidth of the GPUs with the use of APIs that hide the GPU hardware from programmers. This paper presents experimental results on the parallel processing for some well known on-line string matching algorithms using one such GPU abstraction API, the Compute Unified Device Architecture (CUDA).

Result

In this, both the serial and the parallel implementations were compared in terms of running time for different reference sequences, pattern sizes and number of threads. It was shown that the parallel implementation of the algorithms was up to 24x faster than the serial implementation, especially when larger text and smaller pattern sizes were used. The performance achieved is close to the one reported for

similar string matching algorithms. In addition, it was discussed that in order to achieve peak performance on a GPU, the hardware must be as utilized as possible and the shared memory should be used to take advantage of its very low latency. Future research in the area of string matching and GPGPU parallel processing could focus on the performance study of the parallel implementation of additional categories of string matching algorithms, including approximate and two dimensional string matching.

Thierry Lecroq[19] propose a very fast new family of string matching algorithms based on hashing q-grams. The new algorithms are the fastest on many cases, in particular, on small size alphabets. The string matching problem consists in finding one or more usually all the occurrences of a pattern $x = x[0..m - 1]$ of length m in a text $y = y[0..n - 1]$ of length n . It can occur, for instance, in information retrieval, bibliographic search and molecular biology.

Result

In this article they presented simple and though very fast adaptations and implementations of the Wu–Manber exact multiple string matching algorithm to the case of exact single string matching algorithm. Experimental results show that the new algorithms are very fast for short patterns on small size alphabets comparing to the well known fast algorithms using bitwise techniques. The new algorithms are also fast on long patterns (length 32 to 256) comparing to algorithms using an indexing structure for the reverse pattern (namely the Backward Oracle Matching algorithm). This new type of algorithm can serve as filters for finding seeds when computing approximate string matching.

Derek Pao, et.al [20] proposes that a memory-efficient hardware string searching engine for antivirus applications is presented. The proposed QSV method is based on quick sampling of the input stream against fixed-length pattern prefixes, and on-demand verification of variable-length pattern suffixes. Patterns handled by the QSV method are required to have at least 16 bytes, and possess distinct 16-byte prefixes. The latter requirement can be fulfilled by a preprocessing procedure. The search engine uses the pipelined Aho-Corasick (P-AC) architecture developed by the first author to process 4 to 15-byte short patterns and a small number of exception cases. Our design was evaluated using the Clam AV virus database having 82,888 strings with a total size that exceeds 8 MB. In terms of byte count, 99.3 percent of the pattern set is handled by the QSV method and 0.7 percent of the pattern set is handled by P-AC. A pattern with distinct 16-byte prefix only occupies up to three lookup table entries in QSV. The overall memory cost of our system is about 1.4 MB, i.e., 1.4 bit per character of the ClamAV pattern set. The proposed method is memory-based, hence, updates to the pattern set can be

accommodated by modifying the contents of the lookup tables without reconfiguring the hardware circuits.

Hassan Ghasemzadeh[21] proposes that Mobile sensor-based systems are emerging as promising platforms for healthcare monitoring. An important goal of these systems is to extract physiological information about the subject wearing the network. Such information can be used for life logging, quality of life measures, fall detection, extraction of contextual information, and many other applications. Data collected by these sensor nodes are overwhelming, and hence, an efficient data processing technique is essential.

Result

Results show the effectiveness of this approach, both for reliable movement classification and reduction of communication.

HyunJin Kim and Seung-Woo Lee [22] propose a memory-based parallel string matching engine using the compressed state transitions. In the finite-state machines of each string matcher, the pointers for representing the existence of state transitions are compressed. In addition, the bit fields for storing state transitions can be shared. Therefore, the total memory requirement can be minimized by reducing the memory size for storing state transitions

Result

This letter proposed a memory-efficient parallel string matching engine in DFA-based string matching. The proposed string matcher can reduce the memory size for storing the existence of state transitions. In addition, the memory requirements can be reduced by sharing state transitions in the transition table. Considering the experiment results, it is evident that the proposed architecture is useful for reducing the storage cost of the DFA-based string matching engine.

Ali Peiravi and Mohammad Javad Rahimzadeh[23] proposes that String matching is a fundamental element of an important category of modern packet processing applications which involve scanning the content flowing through a network for thousands of strings at the line rate. To keep pace with high network speeds, specialized hardware-based solutions are needed which should be efficient enough to maintain scalability in terms of speed and the number of strings. In this paper, a novel architecture based upon a recently proposed data structure called the Bloomier filter is proposed which can successfully support scalability. The Bloomier filter is a Compact data structure for encoding arbitrary functions, and it supports approximate evaluation queries. By eliminating the Bloomier filter's false positives in a space efficient way, a simple yet powerful exact string matching architecture is proposed that can handle several thousand Strings at high rates and is amenable to on-chip realization. The proposed scheme is implemented

in reconfigurable hardware and compare it with existing solutions. The results show that the proposed approach achieves better performance compared to other existing architectures measured in terms of throughput per logic cells per character as a metric.

In this paper, we use parallel algorithms with multicore processors because with multicore processors we can increase the efficiency and the performance.

III. COMPUTING MODEL WITH MULTICORE PROCESSING

As personal computers have become more prevalent and more applications have been designed for them, the end-user has seen the need for a faster, more capable system to keep up. Speedup has been achieved by increasing clock speeds and, more recently, adding multiple processing cores to the same chip. Although chip speed has increased exponentially over the years, that time is ending and manufacturers have shifted toward multicore processing. However, by

increasing the number of cores on a single chip challenges arise with memory and cache coherence as well as communication between the cores. Coherence protocols and interconnection networks have resolved some issues, but until programmers learn to write parallel applications, the full benefit and efficiency of multi core processors will not be attained [24-27].

IV. PROPOSED SYSTEM ARCHITECTURE

a) System Architecture

System Architecture describes “the overall structure of the system and the ways in which the structure provides conceptual integrity”. Architecture is the hierarchical structure of a program components (modules), the manner in which these components interact and the structure of data that are used by that components. The existing string matching system architecture is as shown in Fig 1 and in this the efficiency is not good.

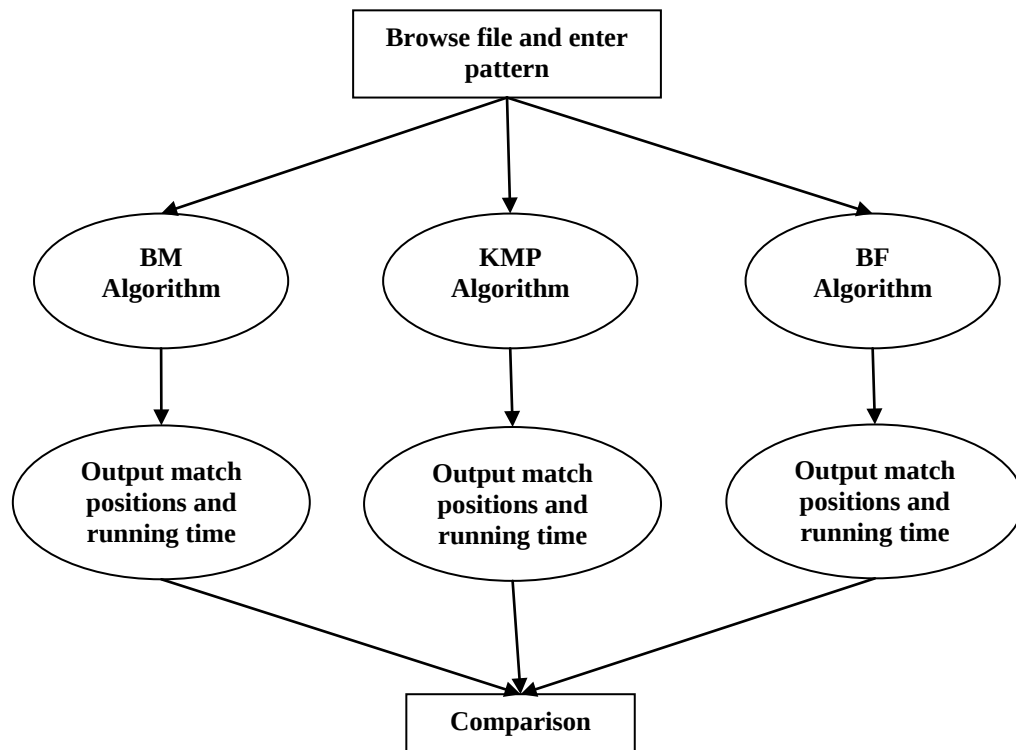


Figure 1 : Existing System

In the existing string matching architecture we search the required pattern sequentially at first we pass the required that is to be searched and this pattern is searched by using the three algorithms Brute force, KMP, Boyer Moore the entire string is passed through all the algorithms and the output match and the running time is calculate for the required pattern from all the algorithms and the algorithm with the least running time is selected, all this is done sequentially which takes more time to execute to improve the efficiency and the

performance in this we use the parallel string matching algorithms with multicore processors as shown in Fig 2.

The proposed system Architecture of Comparison of parallel String Matching Algorithms is as follows in the below diagram. In this search the pattern parallel. in this at first we take the input as a string or text. The required text that is to be searched is further divided into further small patterns and all this patterns are passed on the different parallel algorithms like KMP

boyar Moore, brute force and at all the output position match and running time of all the patterns is calculated and the all the patterns of same algorithm are added and all the resulted running time are compared with

other algorithms resulting time and from them the best one is taken as the efficient algorithm for the string matching.

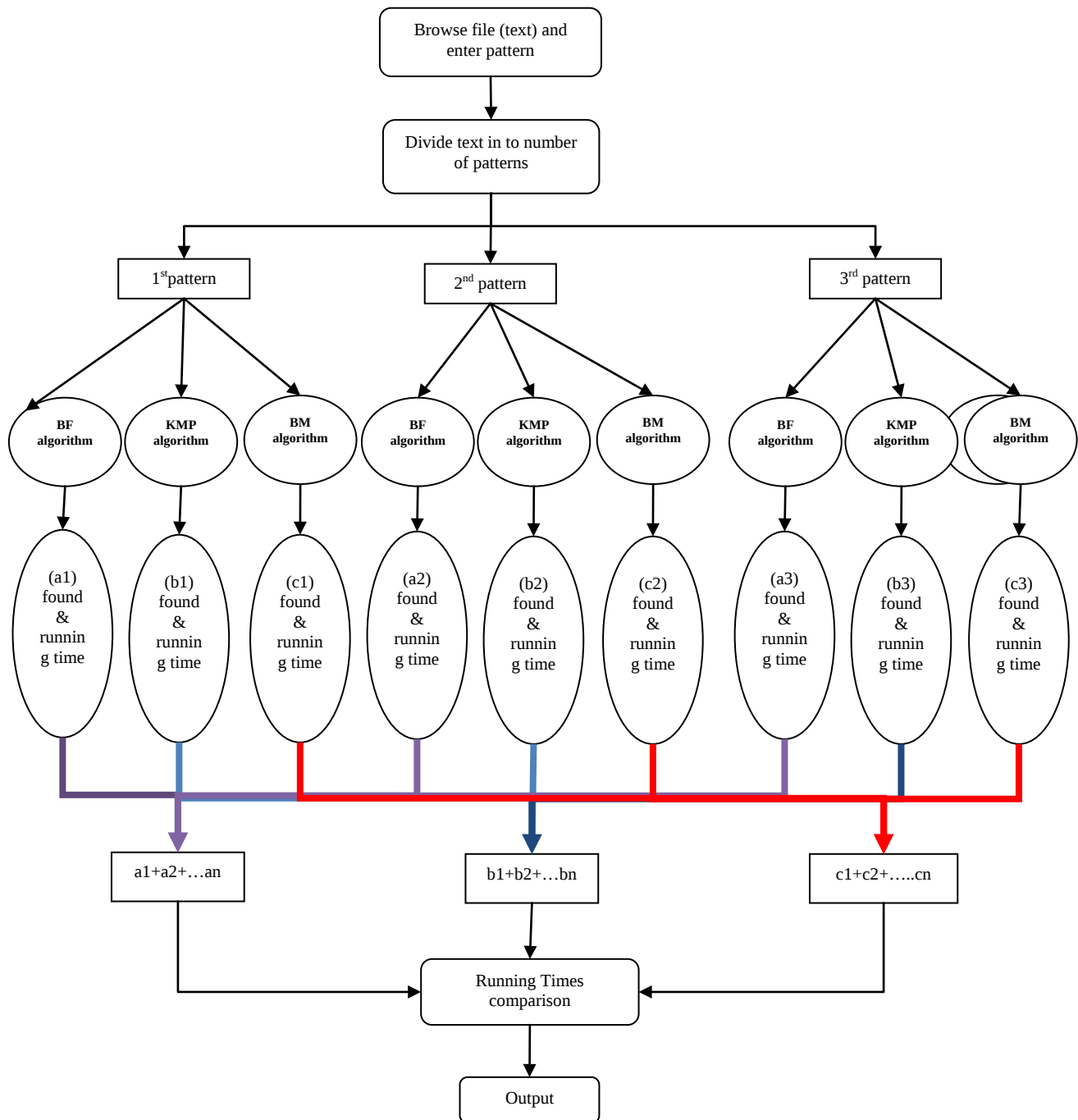


Figure 2 : Proposed Systems

V. PROPOSED APPROACH

In now a day as the current free textual database is growing vast there is a problem of finding the pattern by string matching the efficiency is decreased and takes more time. In our paper, we use parallel algorithms to increase the efficiency on multicore processor we pass the same string to all the

three algorithms and we select the best based on the running time.

a) Implementation

Here we have to implement the proposed system with JAVA 1.7 multi threading, initially we have to implement the BF, KMP, and BM sequentially and then go for parallel implementation with threading on Multicore processor. Here we discuss some of them.

i. *Brute force Algorithm (BF) description and Implementation with parallel programming[28-30]*

The brute force algorithm consists in checking, at all positions in the text between 0 and $n - m$, whether an occurrence of the pattern starts there or not. Then, after each attempt, it shifts the pattern by exactly one position to the right. The brute force algorithm requires no preprocessing phase, and a constant extra space in

addition to the pattern and the text. During the searching phase the text character comparisons can be done in any order. The algorithm can be designed to stop on either the first occurrence of the pattern, or upon reaching the end of the text. This code was run parallel in multiple threads to achieve good efficiency searching, which is shown in Table 1.

Table 1 : Pseudo code for BF

```
FileInputStream fstream = new FileInputStream("F:/multi/genesequence.txt");
DataInputStream in= new DataInputStream(fstream);
BufferedReader br = new BufferedReader(new InputStreamReader(in));
time = System.currentTimeMillis();
while ( ((str = br.readLine()) != null)&&(i<=i1) ){
BruteForceSearch bfs = new BruteForceSearch();
String pattern = "AAGG";
bfs.setString(str, pattern);
first_occur_position = bfs.search();
System.out.println("The text " + pattern + " is first found after the " + first_occur_position + " position.");
i++;}
time = System.currentTimeMillis() - time;
System.out.println("Time elapsed"+time);
```

ii. *Knuth Morris Pratt description and Implementation with parallel programming [28-30]*

Consider an attempt at a left position j , that is when the window is positioned on the text factor $y[j \dots j+m-1]$. Assume that the first mismatch occurs between $x[i]$ and $y[i+j]$ with $0 < i < m$. Then, $x[0 \dots i-1] = y[j \dots i+j-1] = u$ and $a = x[i] \neq y[i+j] = b$. When shifting, it is reasonable to expect that a prefix v of

the pattern matches some suffix of the portion u of the text. Moreover, if we want to avoid another immediate mismatch, the character following the prefix v in the pattern must be different from a . The longest such prefix v is called the *tagged border* of u . This code was run parallel in multiple threads to achieve good efficiency searching, which is shown in Table 2.

Table 2 : Pseudo code for KMP

```
FileInputStream fstream = new FileInputStream("F:/multi/genesequence.txt");
DataInputStream in = new DataInputStream(fstream);
BufferedReader br = new BufferedReader(new InputStreamReader(in));
time = System.currentTimeMillis();
while ( ((str = br.readLine()) != null) && (i<=i1) ) {
KMPS kmp = new KMPS();
String pattern = "AAGG";
kmp.setString(str, pattern);
first_occur_position = kmp.search();
System.out.println("The text " + pattern + " is first found after the " + first_occur_position + " position.");
i++;}
time = System.currentTimeMillis() - time;
System.out.println("Time elapsed"+time);
```

iii. *Boyer Moore Algorithm description and Implementation with parallel programming[28-30]*

The algorithm scans the characters of the pattern from right to left beginning with the rightmost one. In case of a mismatch (or a complete match of the whole pattern) it uses two precomputed functions to shift the window to the right. These two shift functions are called the *good-suffix shift* (also called matching shift) and the *bad-character shift* (also called the

occurrence shift). This code was run parallel in multiple threads to achieve good efficiency searching, which is shown in Table 3.

Table 3 : Pseudo code for BM

```

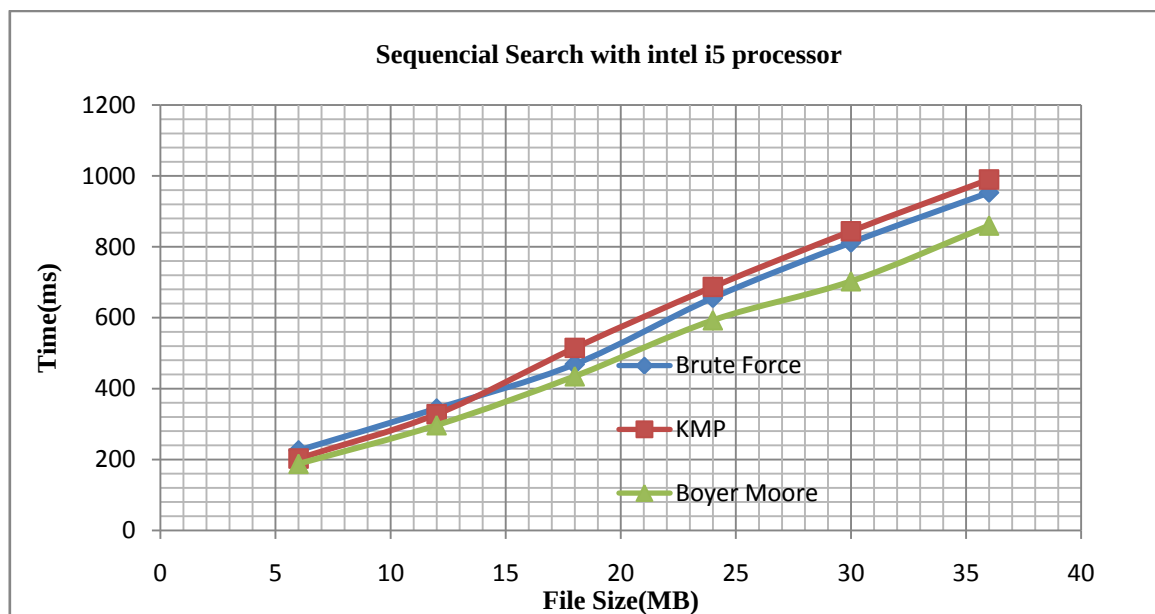
FileInputStream fstream = new FileInputStream("F:/multi/genesequence.txt");
DataInputStream in = new DataInputStream(fstream);
BufferedReader br = new BufferedReader(new InputStreamReader(in));
time = System.currentTimeMillis();
while ( ((str = br.readLine()) != null) && (i<=i1) ){
    BoyerMoore bms = new BoyerMoore();
    String pattern = "AAGG";
    bms.setString(str, pattern);
    first_occur_position = bms.search();
    System.out.println("The text " + pattern + " is first found after the " + first_occur_position + " position.");
    i++;}
time = System.currentTimeMillis() - time;
System.out.println("Time elapsed-in thread-1"+time);

```

b) Claims

Implementation is the stage where the theoretical design is turned into a working system. The most crucial stage in achieving a new successful system and in giving confidence on the system for the users that will work efficiently and effectively. The system will be implemented only after thorough testing and if it is found to work according to the specification. For testing our proposed system we will take the gene sequence data set, consists of the four nucleotides a, c, g and t (standing for adenine, cytosine, guanine, and thymine, respectively) used to encode DNA. Therefore, the

alphabet is $O=\{A, C, G, T\}$. The text is consisted of 7,50,000 records. Our test tested with different processors like i3, i5 etc., here we put some achievements what we develop and observe, finally our system shows that parallel approach is much better than sequential approach with multi core processor. The Fig 3 shows(Graph) Execution time vs File size on sequential search with intel i5 processor using Boyer Moore, Brute force, KMP Algorithm. From the graph we clearly observe that BM is better compared to other approaches.

*Figure 3* : Graph for sequential search (Time vs File size)

The Fig 4 shows(Graph) Execution time vs File size on parallel search with intel i5 processor using Boyer Moore, Brute force, KMP Algorithm. This graph shows the performance difference between Boyer Moore, Knuth Morris Pratt and Brute force algorithms. From the graph we clearly observe that BM is better compared to other approaches. The Fig 5 shows(Graph) Execution time vs File size using Brute force Algorithm on parallel search and sequential search

with intel i5 processor. From the graph we clearly observe that parallel is much better than sequential search.

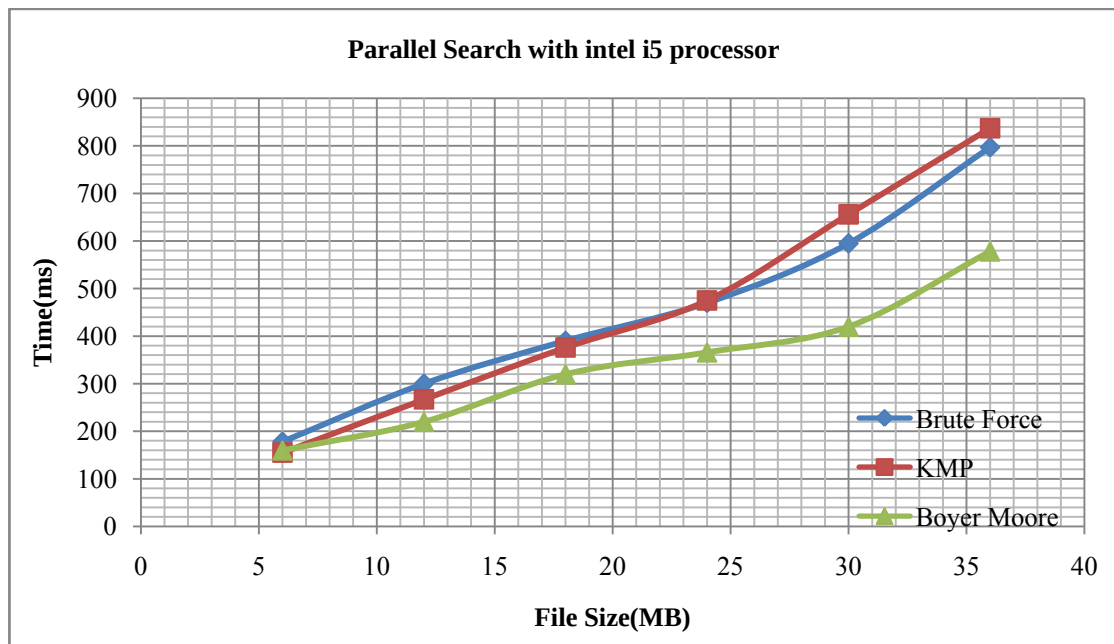


Figure 4 : Graph for parallel search (Time vs File size)

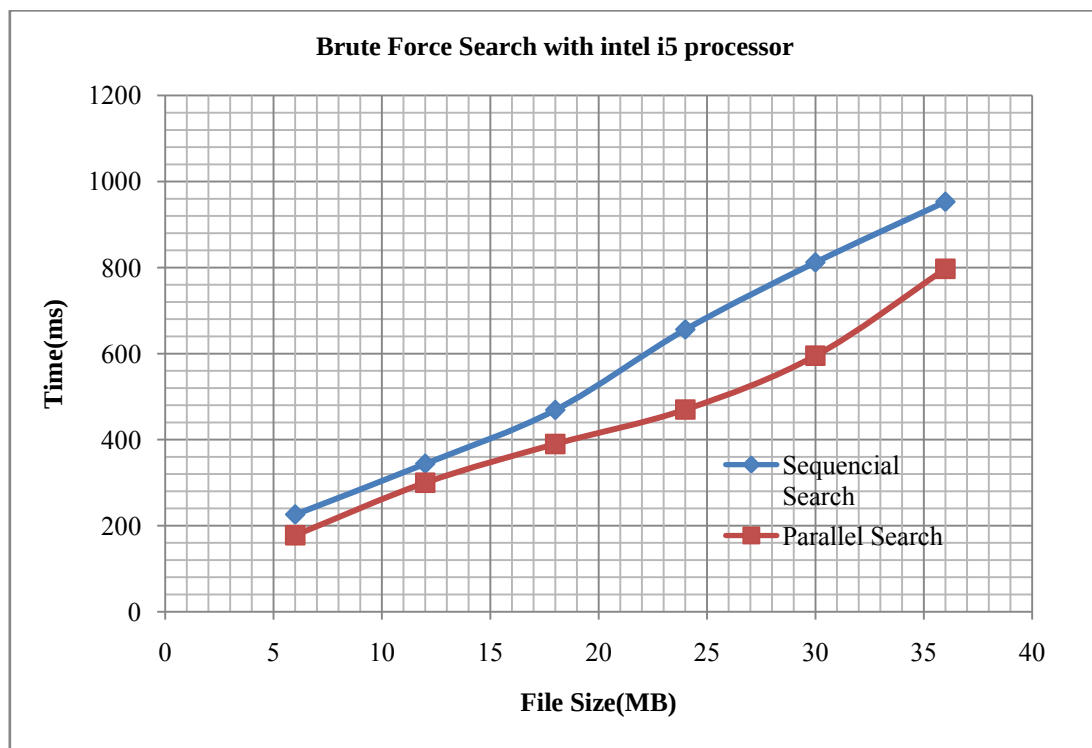


Figure 5 : Graph for Brute force (Time vs File size)

The Fig 6 shows(Graph) Execution time vs File size using KMP Algorithms by sequential search and parallel search with intel i5 processor. From the graph we clearly observe that parallel is much better than sequential search. The Fig 7 shows(Graph) Execution time vs Text length using Boyer Moore Algorithm by parallel search and sequential search with intel i5 processor. From the graph we clearly observe that parallel is much better than sequential search.

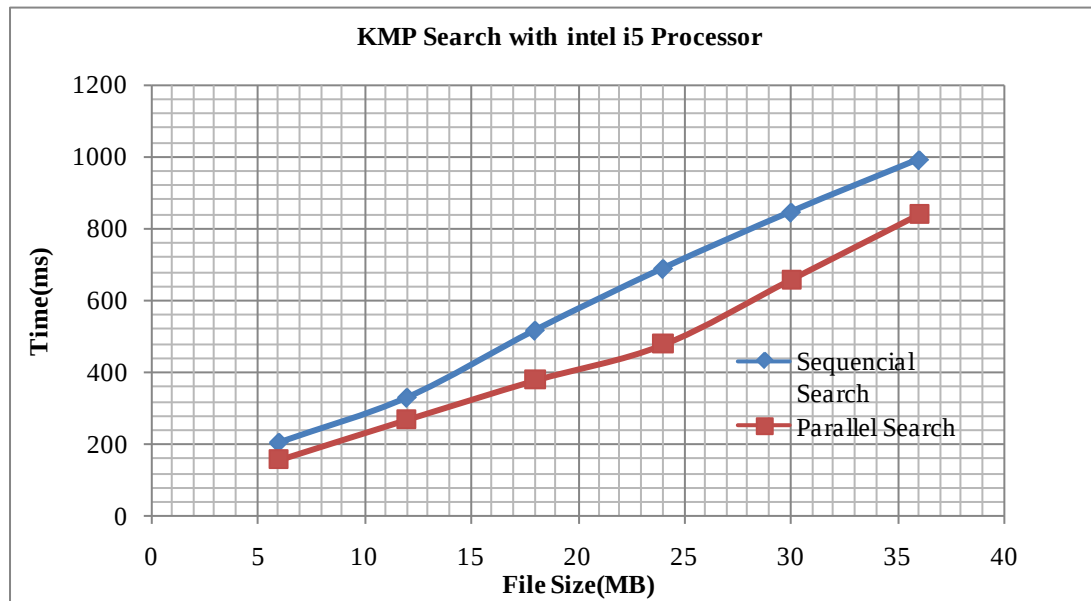


Figure 6 : Graph for KMP (Time vs File size)

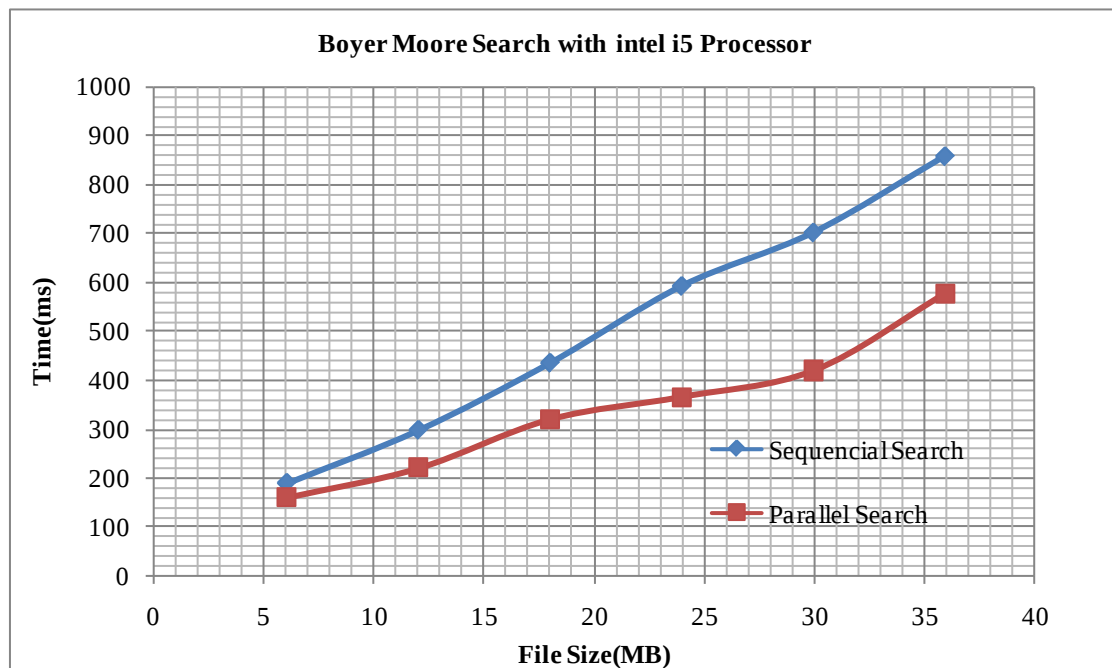
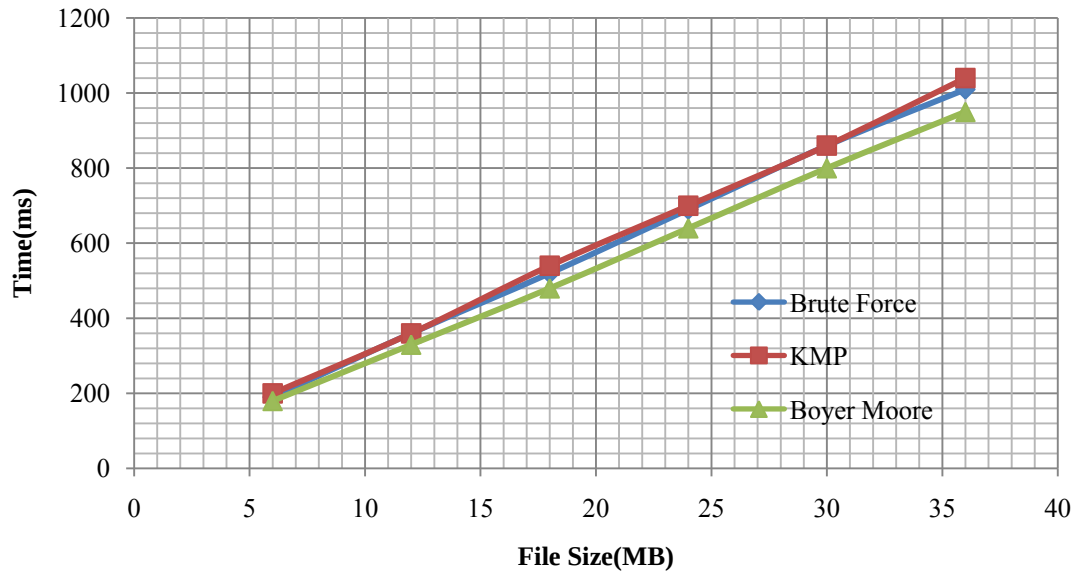
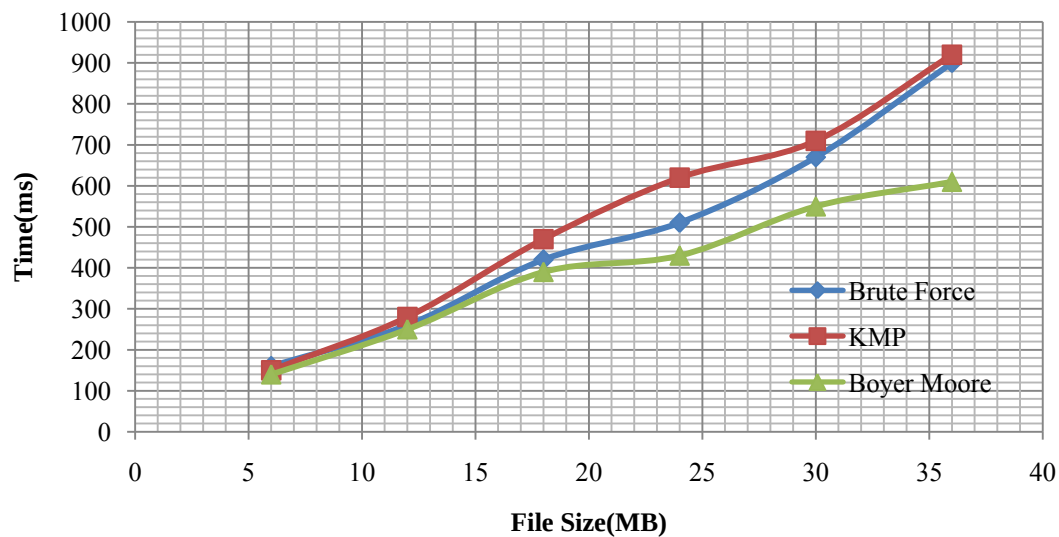


Figure 7 : Graph for Boyer Moore (Time vs Text Length)

The Fig 8 shows (Graph) Execution time vs File size on sequential search with intel i3 processor using Boyer Moore, Brute force, KMP Algorithm. This graph shows the performance difference between Boyer Moore, Knuth Morris Pratt and Brute force algorithms. From the graph clearly observe that BM is better compared to other approaches. The Fig 9 shows (Graph) Execution time vs File size on parallel search with intel i3 processor using Boyer Moore, Brute force, KMP Algorithm. This graph shows the performance difference between Boyer Moore, Knuth Morris Pratt and Brute force algorithms. From the graph

clearly observe that BM is better compared to other approaches, as well as this parallel approach is much better compared to sequential approaches.

Sequential Search with intel i3 Processor*Figure 8* : Graph for sequential search (Time vs File size)**Paralle Search with intel i3 Processor***Figure 9* : Graph for parallel search (Time vs File size)

The Fig. 10 shows (Graph) Execution time vs file size of Brute force sequential search algorithm in Intel i3 and Intel i5 processor. From the graph we says that i5 is performed well compared to i3. The Fig 11 shows (Graph) shows Execution time vs File size using Brute force parallel search algorithm on intel i3 and i5 processor. From the graph we says that i5 is performed well compared to i3.

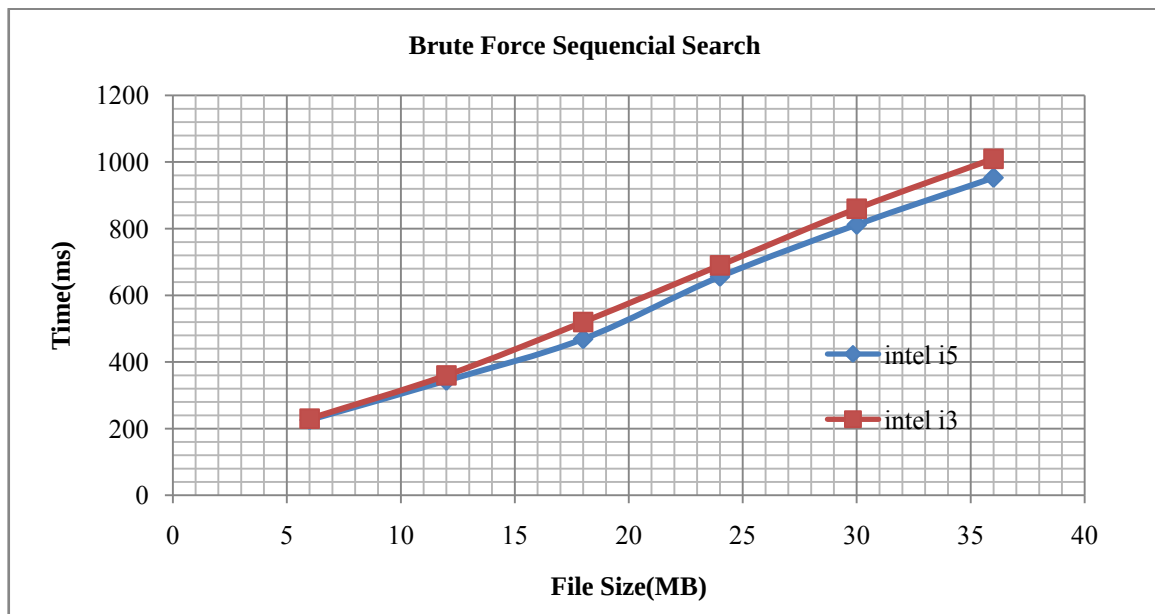


Figure 10 : Graph for Brute force(Time vs file size)

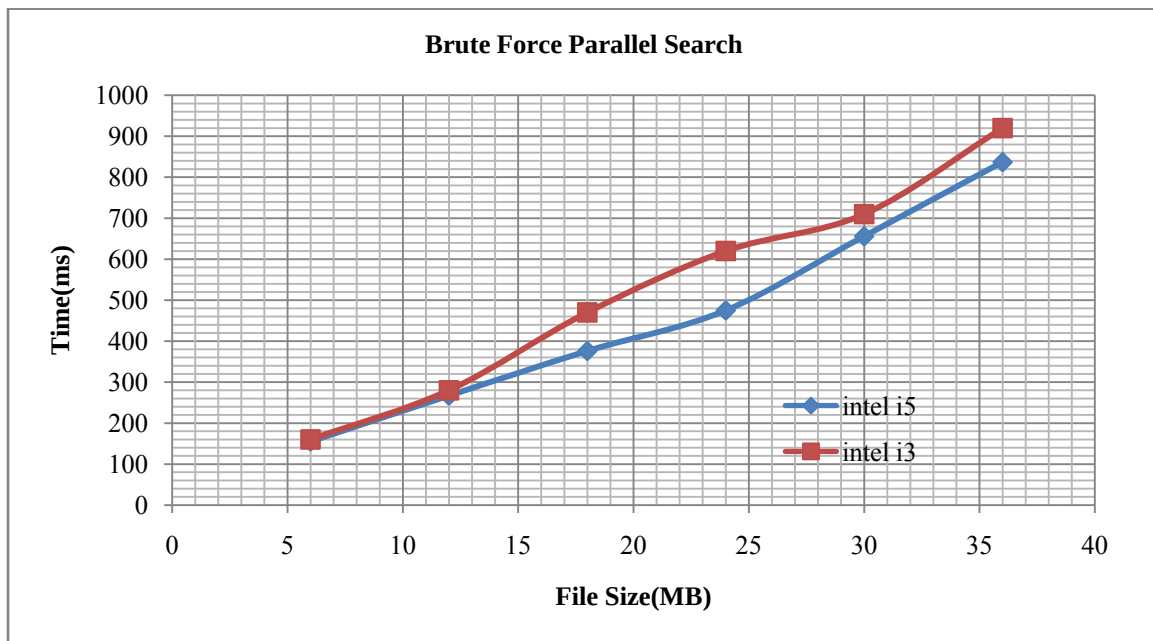


Figure 11 : Graph for Brute force parallel search (Time vs file size)

The Fig 12 shows (Graph) Execution time vs File size using KMP sequential search algorithm on Intel i3 and i5 processor. From the graph we says that i5 is performed well compared to i3. The Fig13 shows (Graph) Execution time vs File size using KMP parallel search algorithm on Intel i3 and i5 processor. From the graph we says that i5 is performed well compared to i3.

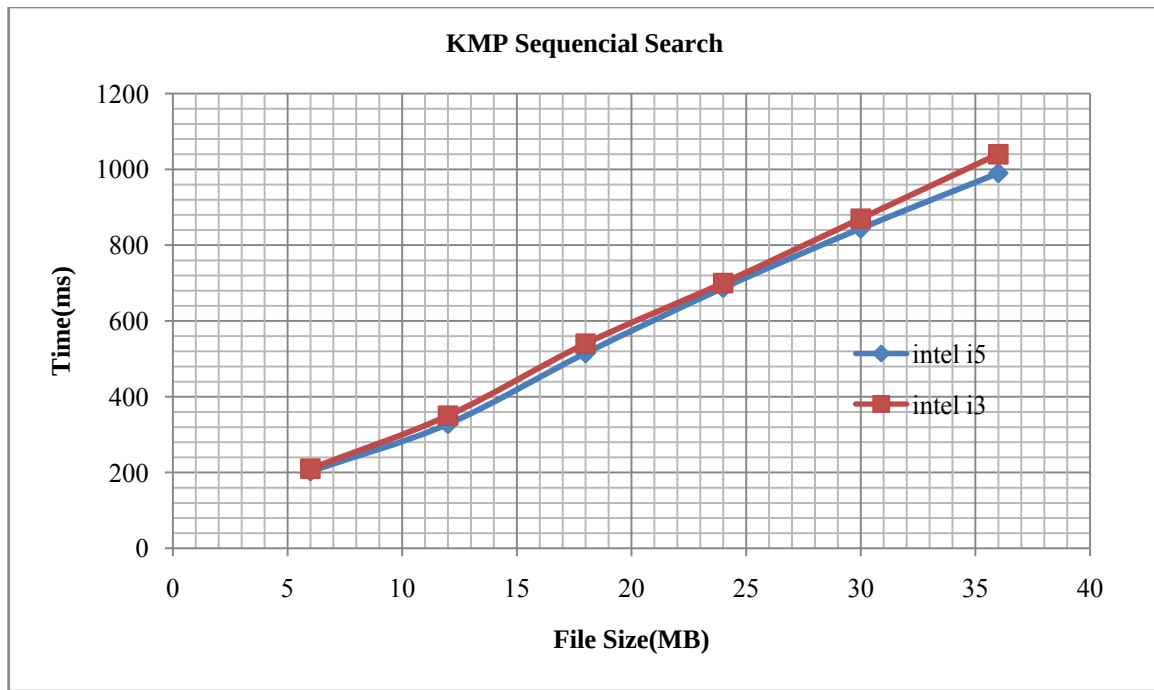


Figure 12 : Graph for KMP Sequential Search (Time vs file size)

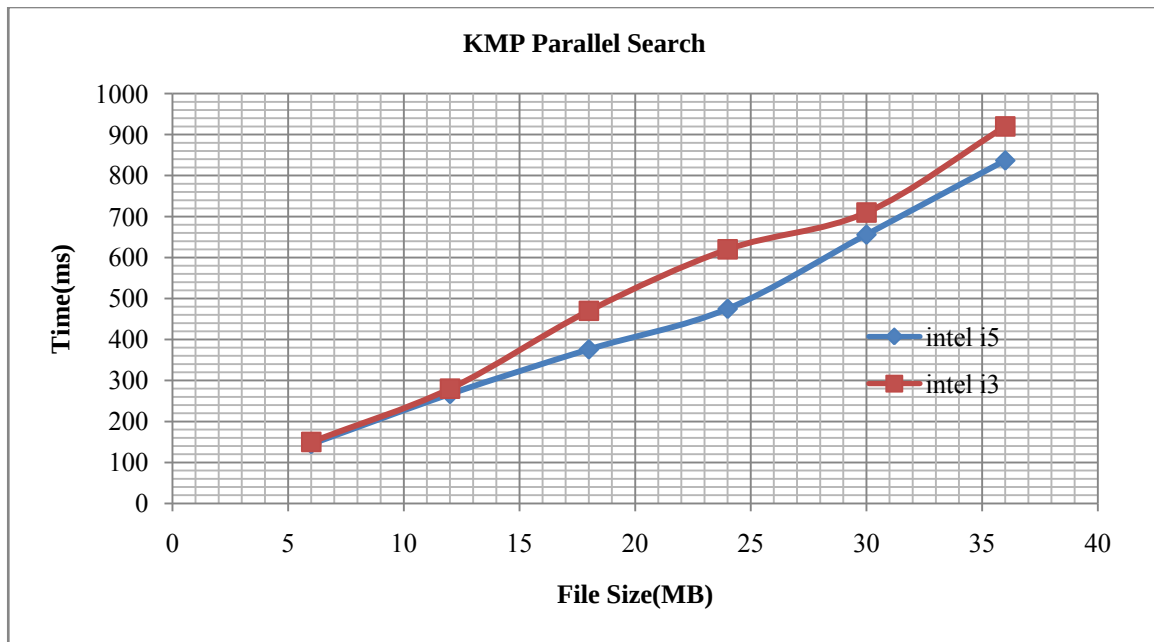


Figure 13 : Graph for KMP parallel Search (Time vs file size)

The Fig 14(Graph) shows Execution time vs File size using Boyer Moore sequential search algorithm on Intel i3 and i5 processor. From the graph we says that i5 is performed well compared to i3. The Fig 15(Graph) shows Execution time vs File size using Boyer Moore parallel search algorithm on Intel i3 and i5 processor. From the graph we says that i5 is performed well compared to i3.

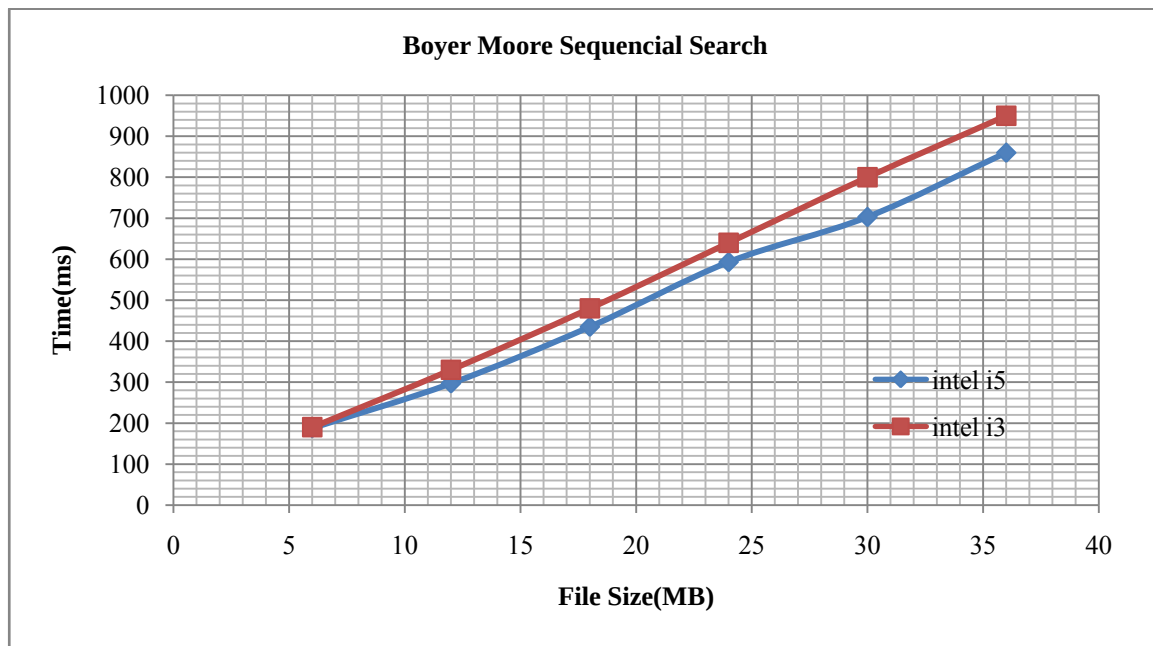


Figure 14 : Graph for Boyer moore sequential Search (Time vs file size)

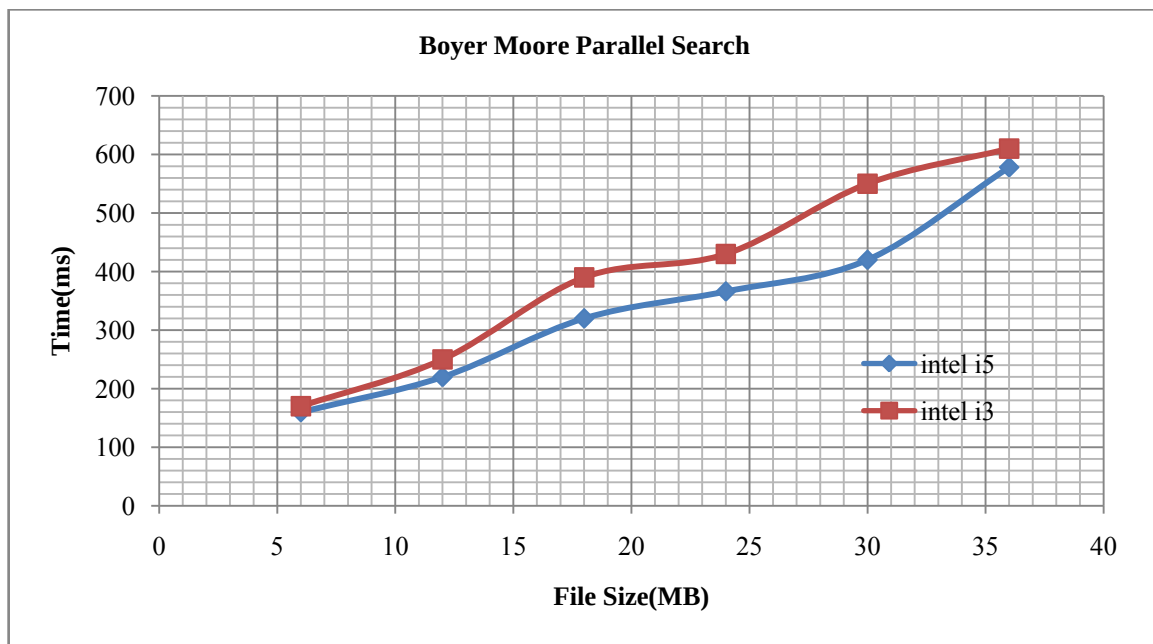


Figure 15 : Graph for Boyer Moore parallel Search (Time vs file size)

VI. CONCLUSIONS

In this paper we performed a comparative study on Knuth Morris Pratt, Boyer Moore and Brute force string matching algorithms based on the running time and in our tests with multicore processing, we used strings of varying lengths and texts of varying lengths. From the test results it is shown that the Boyer Moore algorithm is extremely efficient in most cases and Knuth-Morris-Pratt algorithm is not better on the average than the Brute force algorithm. We conclude that Boyer Moore string matching algorithm is the most efficient

one among the three string matching algorithms with multicore processing compared to earlier versions. As a future enhancement, these algorithms can be compared with other efficient parallel string matching algorithms thereby finding the most efficient algorithm which can be used in many fields such as cryptography, molecular biology. Thus the problem of matching becomes easier.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Badoiu M. and Indyk P., "Fast Approximate Pattern Matching with Few Indels via Embeddings," ACM-

- SIAM Symposium on Discrete Algorithms, pp. 651-652, 2004.
2. H. Kim, H. Hong, H.-S. Kim, and S. Kang, "A Memory-Efficient Parallel String Matching for Intrusion Detection Systems," IEEE Comm. Letters, vol. 13, no. 12, pp. 1004-1006, Dec. 2009.
3. Hyun Jin Kim, Hong-Sik Kim, Sungho Kang, "A Memory-Efficient Bit-Split Parallel String Matching Using Pattern Dividing for Intrusion Detection Systems," IEEE Transactions on Parallel and Distributed Systems, vol. 22, no. 11, pp. 1904-1911, Nov. 2011.
4. Mhashi M., Rawashdeh A., and Hammouri A., "A Fast Approximate String Searching Algorithm," Computer Journal of Science Publication, pp. 405-412, 2005.
5. Dany Breslauer, Leszek Gasieniec, Roberto Grossi, "Constant-Time word-size string matching", Proceedings of the 23rd Annual conference on Combinatorial Pattern Matching, Helsinki, Finland, 2012.
6. S. Faro and T. Lecroq, "The Exact Online String Matching Problem: a Review of the Most Recent Results", ACM Computing Surveys 45(2), 2013.
7. Jonathan L., "Analysis of Fundamental Exact and Inexact Pattern Matching Algorithms," Technical Document, Stanford University, 2004.
8. 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.
9. 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.
10. 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.
11. 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.
12. 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.
13. 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.
14. S. viswanadha raju, "parallel string matching algorithm using grid", "international journal of distributed and parallel systems" (ijdps) vol.3, no.3, 2012.
15. Hyunjin kim and sungho kang, "A pattern group partitioning for parallel string matching using a pattern grouping metric", ieee communications letters, vol. 14, no. 9, pp. 878-880, 2010.
16. Daniel luchaup, I "speculative parallel pattern matching", "ieee transactions on information forensics and security, vol. 6, no. 2, pp.438-451, 2011".
17. Hyunjinkim, member, ieee, "A memory-efficient bit-split parallel string matching using pattern dividing for intrusion detection systems", ieee transactions on parallel and distributed systems, vol. 22, no. 11, pp.1904-1911, 2011".
18. Charalampos S, "String matching on a multicore gpu using cuda", "parallel and distributed processing laboratory department of applied informatics, university of macedonia 156 egnatia str, p.o. box 1591, 54006 thessaloniki, greece".
19. Thierry lecroq, "Fast exact string matching algorithms", "litis, faculty des sciences et des techniques, university de rouen, 76821 mont-saint-avignon cedex, france", january 2007.
20. Derek pao, "string searching engine for virus scanning" ieee transactions on computers, vol. 60, no.11, pp1596-1609, 2011.
21. Hassan ghasemzadeh, "structural action recognition in body sensor networks: distributed classification based on string matching", "ieee transactions on information technology in biomedicine, vol. 14, no. 2, pp.425-435, 2010".
22. HyunJin Kim and Seung-Woo Lee, "A Hardware-Based String Matching Using State Transition Compression for Deep Packet Inspection", "ETRI Journal, Volume 35, Number 1, pp-154-157, 2013"
23. Ali Peiravi and Mohammad Javed Rahimzadeh, "A novel scalable and storage-efficient architecture for high speed exact string matching", "etri journal, volume 31, number 5, pp.545-553,2009".
24. Geist, A., Beguelin, A., Dongarra, J., Jiang, W., Manchek, R. & V., S. "PVM: Parallel Virtual Machine", A Users Guide and Tutorial for Networked Parallel Computing, MIT Press. 1994.
25. Grama, A., Karypis, G., Kumar, V. & Gupta, A. "Introduction to Parallel Computing", Addison Wesley, 2003.
26. Leow, Y., Ng, C. &W.F.,W.. "Generating hardware from OpenMP programs", International Conference on Field Programmable Technology, pp. 73–80, 2006.
27. Marr, D. T., Binns, F., Hill, D. L., Hinton, G., Koufaty, D. A., Miller, J. A. & Upton, M.. "Hyper-Threading

- Technology Architecture and Microarchitecture", Intel Technology Journal, pp. 4–15, 2002.
28. R. S. Boyer and J. S. Moore, A fast string searching algorithm, Carom. Association of computing Machinery, pp-262–272, 1977.
 29. Donald Knuth; James H. Morris, Jr, Vaughan Pratt, "Fast pattern matching in strings". SIAM Journal on Computing, pp-323–350, 1977.
 30. Z. Galil, "Optimal parallel algorithms for string matching," in Proc. 16th Annu. ACM symposium on Theory of computing, pp. 240-248, 1984.





This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 1 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Analysis of Parallel Boyer-Moore String Search Algorithm

By Abdullellah A. Alsaheel, Abdullah H. Alqahtani & Abdulatif M. Alabdulatif

King Saud University, Saudi Arabia

Abstract - Boyer Moore string matching algorithm is one of the famous algorithms used in string search algorithms. Widely, it is used in sequential form which presents good performance. In this paper a parallel implementation of Boyer Moore algorithm is proposed and evaluated. Experimental results show that it is valuable with zero overhead and cost optimal. The comparison between sequential and parallel showed that the parallel implementation was faster and more useful.

Keywords : *parallel processing, parallel algorithm, boyer-moore, string matching algorithm, performance analysis.*

GJCST-A Classification : *F.2.0*



Strictly as per the compliance and regulations of:



Analysis of Parallel Boyer-Moore String Search Algorithm

Abdullellah A. Alsaheel ^α, Abdullah H. Alqahtani ^σ & Abdulatif M. Alabdulatif ^ρ

Abstract - Boyer Moore string matching algorithm is one of the famous algorithms used in string search algorithms. Widely, it is used in sequential form which presents good performance. In this paper a parallel implementation of Boyer Moore algorithm is proposed and evaluated. Experimental results show that it is valuable with zero overhead and cost optimal. The comparison between sequential and parallel showed that the parallel implementation was faster and more useful.

Keywords : parallel processing, parallel algorithm, boyer-moore, string matching algorithm, performance analysis.

I. INTRODUCTION

String matching algorithm is a huge area used when there is a need to search for the occurrence of a pattern of a string in a text. The applications of string matching algorithms is widely used like in

security field when a protection system searching for a malicious pattern in a big text string. Boyer Moore is one of the most algorithms used in this field. The Boyer Moore algorithm search for the occurrence of pattern $P[1-N]$ in a text $T[1-M]$ by comparing the pattern in the text from right to left and shifting the pattern from left to right when mismatch occur using one of two functions that are bad character shift or good suffix shift [1].

The sequential Boyer Moore searching algorithm requires $O(NM)$ in worst case when the pattern occurs in the text. And in the best performance is $O(M/N)$ [4], where length of a pattern is N , and length of a text string is M . Actually, This algorithm is not efficiently working with small size strings but it is efficient with big size strings.

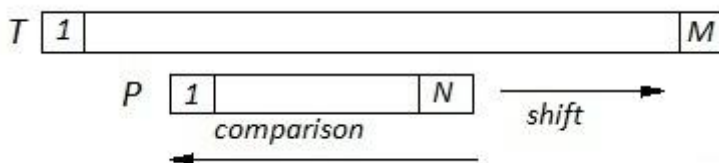


Figure 1 : BM algorithm mechanism make comparison from right to left and shift from left to right when mismatch occur

In this paper implementing a parallel Boyer Moore is proposed and evaluating its performance and compare it with sequential version. Most of performance metrics used in evaluating any algorithm is used to make a clear comparison between sequential and parallel algorithm and presents valuable metrics for algorithm.

II. RELATED WORKS

A few works on parallelizing Boyer Moore algorithm proposed. In [2] implements a parallel algorithm in clustered computing environment for many kinds of string matching algorithms and Boyer Moore algorithm is one of them. They made a comparison between the performance of this algorithm in sequential and parallel on different sizes of data and in different number of nodes. The result for their experiment and comparison demonstrated that Boyer Moore in parallel is taken less time than in sequential.

Popa et al [3] has made an implementation for a parallel algorithm for Boyer Moore and calculating the speed up and efficiency in various numbers of processes. He used single program multiple data stream (SPMD) as a model and C as programming language with using parallel virtual machine (PVM) functions.

III. THE PARALLEL ALGORITHM

Parallel processing is the concept of executing the same problem on different processing units concurrently. The parallel implementation is made to work on one control unit with different data that called single instruction stream multiple data stream (SIMD). The communication model is shared address space where the algorithm implemented on the same platform.

In this algorithm different tasks assigned to P processors to present results. While these tasks are homogeneous and there are no dependencies between them, the bag of tasks has been used as a design pattern for the parallel algorithm.

Authors ^{α σ ρ} : King Saud University, College of computer and information sciences, Saudi Arabia. E-mail : cs_saheel@hotmail.com

The idea is to divide the text T to multiple partitions and make the comparison over the same copy

of pattern and every worker concurrently implement M/P [1] where M is the length of text T .

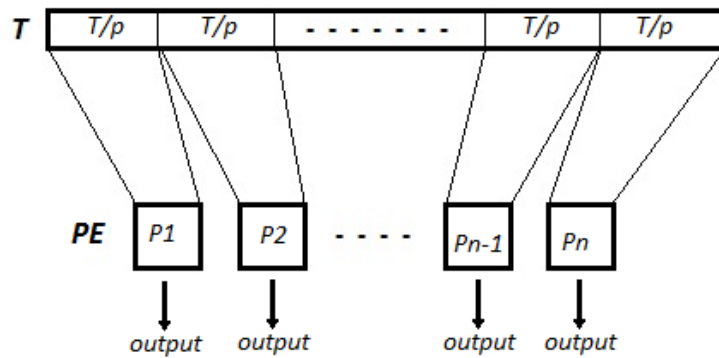


Figure 2 : Partitioning different parts of T on available P

The partitioning is occurred between lines of file text and cannot be occurred in the middle of line and assuming that the pattern that being searched for, will not be occurred between any two lines.

a) Theoretical Analysis

Sequential algorithm has been evaluated by its serial runtime T_s , that taken from the start of the execution of the algorithm to the end of it. The sequential time of Boyer Moore when pattern occur is $T_s = O(MN)$ in worst case [4][5]. The parallel runtime that denoted by T_p is the time taken from the first processor starts to the end of last processor. When our algorithm depends on dividing a text T on the available workers P , the parallel time will be then $T_p = O(\frac{M}{P} * N)$ and the total cost is $pT_p = O(M * N)$.

In the following some of the most famous performance metrics used for parallel systems measurement applied on the proposed parallel algorithm [6]:

- Total overhead: the total time that spent by processors over sequential time $T_o = pT_p - T_s = 0$.
- Speed Up S : the ration of time taken by the algorithm taken in sequential to the same algorithm in parallel, $S = P$.
- Efficiency E : is the relation between speed up and number of processors, $E = 1$ which is in this case is optimal.

By applying these metrics on parallel algorithm, it showed that using Boyer Moore parallel algorithm is more efficient and better than using Boyer Moore sequential algorithm.

b) Practical Analysis

The algorithm has been implemented by Java programming language with having sequential and parallel versions of Boyer-Moore algorithm. The experiment has been applied on a platform with 8 processing units and 8 GB RAM (Random Access Memory).

This implementation used bad rule character that scan the pattern starting from the rightmost character against the text and the pattern will be shifted to the right whenever a mismatch occurs, pattern shifting will be according to the number of positions that stated at skip table that has stored the right most occurrence of that mismatched character at the pattern.

The pseudo-code of Boyer-Moore algorithm:

```
int tableSize;
int[] table;
String pattern;
public BoyerMoore(String pattern) {
    this.tableSize = 256;
    this.pattern = pattern;
    table = new int[tableSize];
    for (int c = 0; c < tableSize; c++)
        table[c] = -1;
    for (int j = 0; j < pattern.length(); j++)
        table[pattern.charAt(j)] = j;
}
public int search(String text) {
    int N = pattern.length();
    int M = text.length();
    int skip;
    for (int i = 0; i <= M - N; i += skip) {
        skip = 0;
        for (int j = N-1; j >= 0; j--) {
            if (pattern.charAt(j) != text.charAt(i+j)) {
                skip = Math.max(1, j - table[text.charAt(i+j)]);
                break;
            }
        }
        if (skip == 0)
            return i;
    }
    return M;
}
```

Partitioning the text into sub-amounts of work will divide the problem into smaller sub-problems, so

then multiple workers can work concurrently at the same time. All of the workers require an access to the

complete list of patterns to compare their portions of text against all patterns.

Number of Elements (N)	Sequential Time (Milliseconds)	Parallel Time (Milliseconds)
500	50	5
1000	96	16
1500	143	27
2000	191	40
2500	232	61
3000	278	82
3500	328	114
4000	367	140
4500	413	145
5000	461	160

Table 1 : Sequential Processing P=1, Parallel Processing P=8

The performance of the sequential algorithm is stable and easy to read as Table1 has demonstrated, since adding 500 elements each time will requires 50 Milliseconds approximately; despite the number of data elements being processed it will usually consume the same amount of time for processing the same amount of data. On the other hand there is the performance of the parallel processing which is presented at Table1 and

may requires more analysis to be understood, however the table demonstrated that adding 500 data elements at each time will affect the processing time significantly whenever N is small; whereas adding the same amount of data to large number of elements will not affect the processing time that much. Conclusion is that the parallel version of Boyer-Moore algorithm works better with larger data.

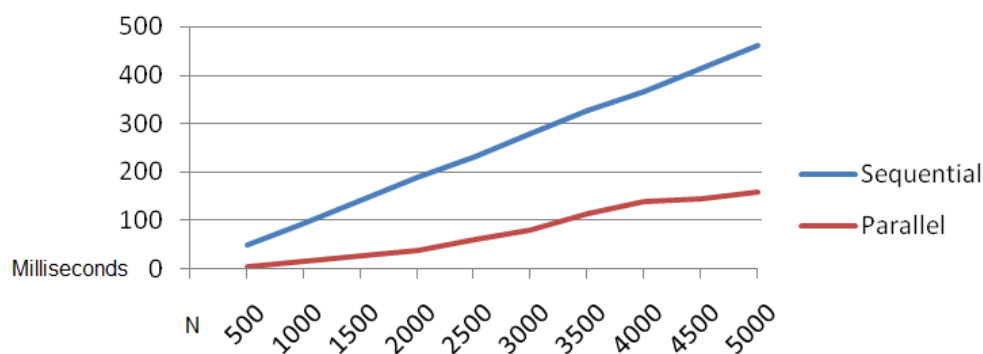


Figure 3 : N is the number of text lines being processed

Number of Processors (P)	Actual Time (Milliseconds)
1	277
2	170
3	134
4	103
5	101
6	106
7	93
8	82

Table 2 : Varying number of Processors, N=3000

Adding more workers will significantly improve the processing time especially whenever the number of workers is small. Notice that adding more workers not necessarily will improve processing performance, this due to the fact that Boyer-Moore algorithm will have worse case whenever there are more mismatches which require more shifting to do; and since the parallel algorithm has partitioned the data across multiple workers, that partitioning sometimes may not guarantee a better performance then what would you have if you have less workers especially whenever the number of workers is large where adding one more worker may not improve the performance significantly. Some data partitioning introduces an additional delay since some

data portions that held by some workers may have more mismatches and more shifting than what other workers may need. So the data nature and the number of workers are critical for the parallel algorithm

performance especially when there is large number of elements to be processed and where adding more workers may not significantly improve the performance.

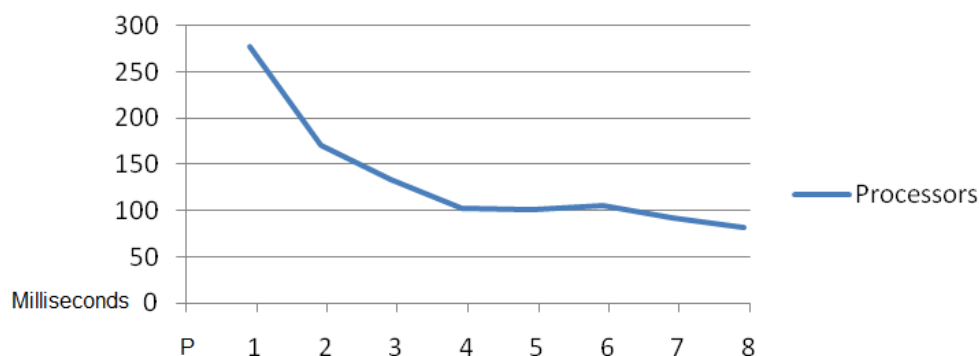


Figure 4 : P is the number of processors

Number of Threads (T)	Actual Time (Milliseconds)
10	61
12	49
14	40
16	38
18	38
20	35
22	35
24	41
26	38
28	29

Table 3 : Parallel processing with varying number of workers, P=8, N=3000

To have more workers than the existing number of processors may lead to significant performance improvement, since having more workers means having smaller sub-problems which are faster to be processed. But one should consider that having more workers than the existing number of processors will introduce an additional overhead due to the context switching. So, identifying a threshold-like which specifies the maximum allowed number of workers that can work safely in performance wise. The experiment showed that on a platform with P=8 processors, there is a safe and better performance can be achieved by having $2 * P = 16$ threads. Increasing the number of threads more than that may result in better or worse performance so one is advised to not exceed that threshold due to instability.

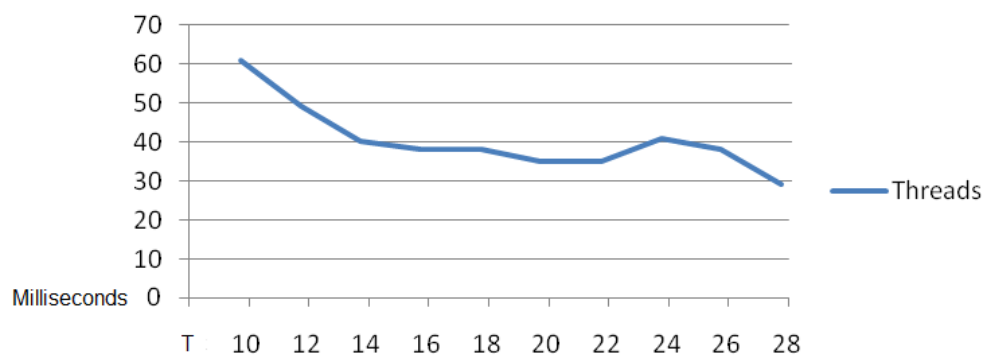


Figure 5 : T is the number of threads

IV. CONCLUSION

Sequential and parallel implementations of Boyer-Moore string matching algorithm have been evaluated. Theoretically, it proved that the performance of the parallel algorithm is cost optimal with zero overhead. In practical experiment, different sizes of data have been used to show that the parallel algorithm is

very faster and better than sequential especially when the data is large. Also, using different number of processors we demonstrated that the parallel Boyer-Moore works better when the number of workers get increased unless the new data partitioning result in having some workers with data has many shifting which end up in delaying all of the workers. Number of workers/threads threshold has been proposed that can

specify the number of threads which can be used efficiently with having stable performance, $2 \cdot P$ threads, where P is the number of existing processors.

REFERENCES RÉFÉRENCES REFERENCIAS

1. R. S. Boyer and J. S. Moore, "A fast string searching algorithm," Communications of the ACM, vol.20, Session10, Oct.1977, pp.761– 772.
2. P. J. C, and K. S. Panicker, "Single Pattern Search Implementations in a Cluster Computing Environment", in 4th IEEE International Conference on Digital Ecosystems and Technologies, 2010.
3. D. Breslauer and Z. Galil, "A Parallel implementation of Boyer-Moore String Searching Algorithm," Sequencess II, 1993, pp. 121-142.
4. Z. Galil, "On improving the worst case running time of the Boyer-Moore string matching algorithm" Communication of the. ACM Vo. 22, no. 9, pp. 505–508, Sept. 1979.
5. R. Cole, "Tight bounds on the complexity of the Boyer-Moore string matching algorithm," proceeding in SODA '91 Proceedings of the second annual ACM-SIAM symposium on Discrete algorithms USA. Philadelphia, pp. 224-233, 1991.
6. A. Grama, G. Karypis, V. Kumar and A. Gupta, "Introduction to parallel computing, second edition," Addison-Wesley, 2003.



This page is intentionally left blank



GLOBAL JOURNAL OF COMPUTER SCIENCE AND TECHNOLOGY
HARDWARE & COMPUTATION
Volume 13 Issue 1 Version 1.0 Year 2013
Type: Double Blind Peer Reviewed International Research Journal
Publisher: Global Journals Inc. (USA)
Online ISSN: 0975-4172 & Print ISSN: 0975-4350

Software Defined 8, 16 & 24 Bit Digital Logic Design by One Microcontroller

By Md. Khaled Hossain & Md. Nasimuzzaman Chowdhury

AIUB-American International University, Bangladesh

Abstract - Now-a-days digital circuits are getting more complex. One IC in a digital circuit is used for a fixed purpose and its operation cannot be defined through software. Because of this limitation digital circuit becomes larger in size. When designing an 8bit digital circuit we do not include 16bit or 24bit components, but this limits our scope of design and versatility of the design. To overcome this problem an 8bit microcontroller is programmed which is able to do addition, subtraction, multiplication and 28 other digital operations in 8, 16 & 24 bit level. To add six 8 bit data 5 adder ICs are not needed anymore. This IC can do it all alone. For any logic operation the regarding mode needs to be selected in the same IC to perform desired operation. It is software defined digital logic design IC. This IC will save time, space & reduce cost in digital circuit designing.

Keywords : *microcontroller, digital logic IC, universal shift register, encoder, decoder.*

GJCST-A Classification : *B.7.1*



Strictly as per the compliance and regulations of:



Software Defined 8, 16 & 24 Bit Digital Logic Design by One Microcontroller

Md. Khaled Hossain^α & Md. Nasimuzzaman Chowdhury^σ

Abstract - Now-a-days digital circuits are getting more complex. One IC in a digital circuit is used for a fixed purpose and its operation cannot be defined through software. Because of this limitation digital circuit becomes larger in size. When designing an 8bit digital circuit we do not include 16bit or 24bit components, but this limits our scope of design and versatility of the design. To overcome this problem an 8bit microcontroller is programmed which is able to do addition, subtraction, multiplication and 28 other digital operations in 8, 16 & 24 bit level. To add six 8 bit data 5 adder ICs are not needed anymore. This IC can do it all alone. For any logic operation the regarding mode needs to be selected in the same IC to perform desired operation. It is software defined digital logic design IC. This IC will save time, space & reduce cost in digital circuit designing.

Keywords : microcontroller, digital logic IC, universal shift register, encoder, decoder.

1. INTRODUCTION

Microcontrollers are capable of executing all digital logic design operations. In another paper we have implemented 28 digital logic design operation by one microcontroller [1]. But still if a 16bit operation is needed we have to cascade two 8bit microcontrollers. Instead of cascading two microcontrollers, it is possible to do operations up to 24bits with this IC. This IC can perform AND, OR, NOR, NAND, XOR, XNOR, Universal Shift Register, Counter operations etc. The Controller used in this project is Atmega1280. It has 86 GPIO. In this project there are 6sets of 8bit input, 25bit output, 6bit main operation selection option and additional 5bit to select sub operations. When user selects 16 bit operation mode, this IC combines Input1 & Input2 as 1st input of 16bit. Combination of Input3 & Input4 becomes 2nd input of 16bit. Input5 & Input6 are combined for 3rd input of 16bit. When user wants to do operation of 24 bit that time Input1, Input2 & Input3 combines & form 1st 24bit input. The rest Input4, Input5 & Input6 act as 2nd input of 24bit. And all the operations output is given through 25bit. Here 25th bit output is the carry bit for 24bit operations.

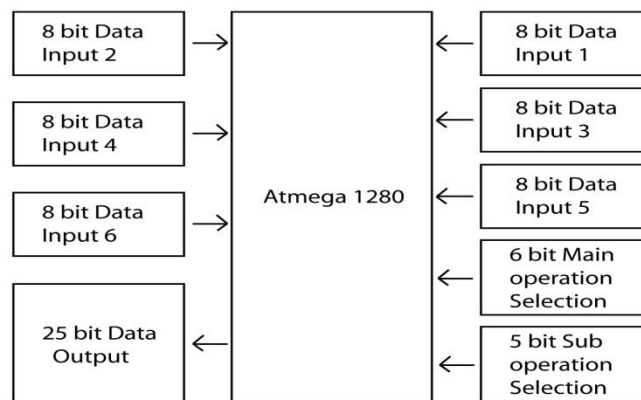


Figure 1 : Basic Project outline

Here input1, input2, input3, input4, input5 & input6 are accordingly portB, portC, PortF, portJ, portK, portL and outputs are portC, portD, portE & portG6_bit. Rest of the pins of portG represent sub operations and to select main operation portH is used. Enable_bar means to enable the IC we need to make portH6_bit low. PortH7 represent synchronous clock input. When 16bit operations are selected PortA are lower eight bits & PortB are higher eight bits. In the same way for input2 & input3 portF, Portk are lower eight bits and PortJ, PortL are higher eight bits accordingly. For 24bit operation in input1 portA are lower eight bits, portB are higher eight bits, portF are most significant eight bits. For second input PortJ are lowest significant bits, portK are higher eight bits, portL are highest eight bits. At output, portC are most significant 8bits, portD are higher 8bits, PortE are more higher 8bits & portG6_bit is the MSB.

Author α : Dept. of Electrical & Electronic Engineering. AIUB-American International University Bangladesh.

E-mails : tokyo_emb@yahoo.com, Mob. +8801944665975

Author σ : Dept. of Electrical & Electronic Engineering. AIUB-American International University Bangladesh.

E-mail : chotonme@yahoo.com, Mob. +8801723209026

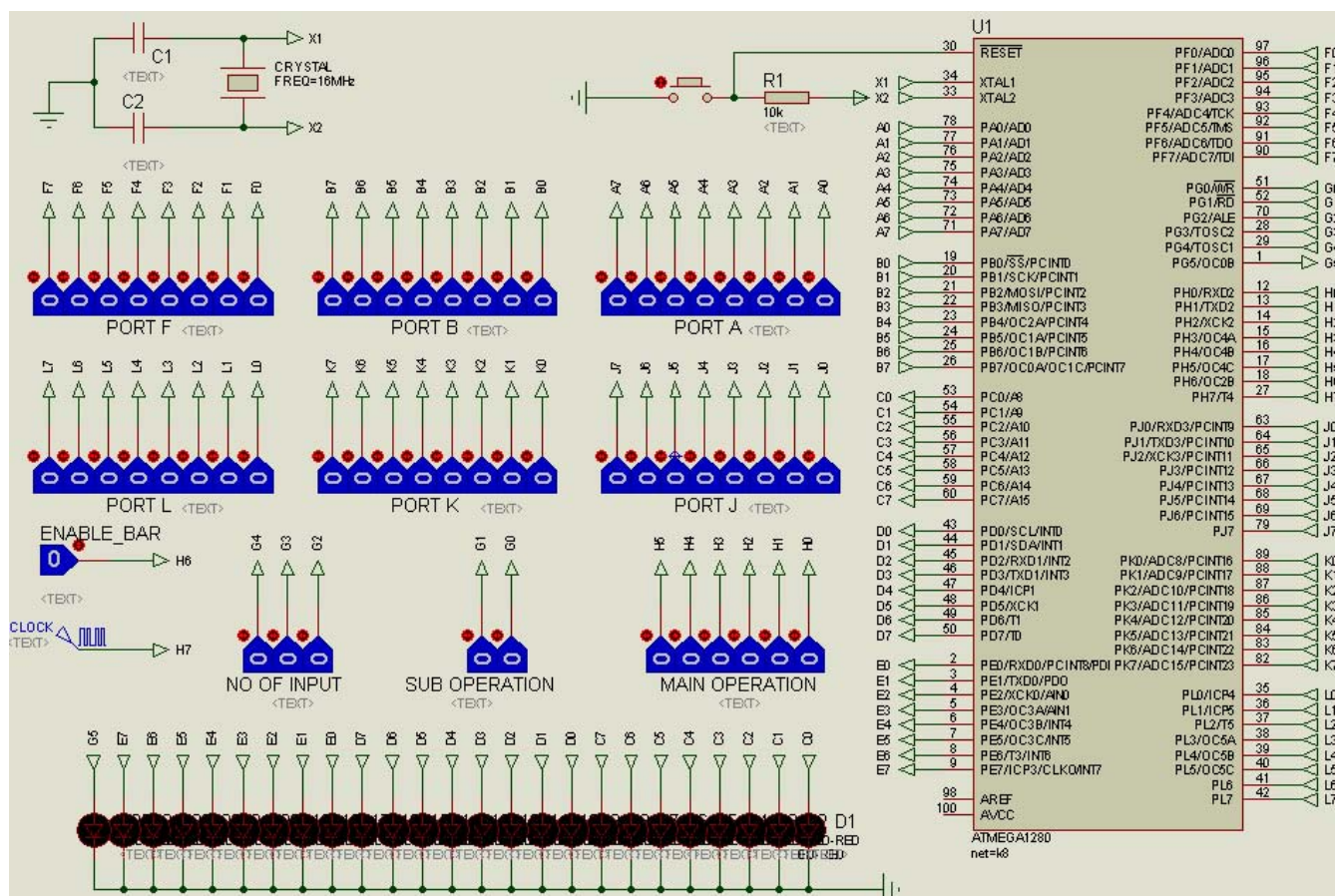


Figure 2 : Main Circuit Diagram

II. CIRCUIT DESCRIPTION

Atmega1280 is used as microcontroller in this project. Operating voltage of atmega1280 is +5v. Here VCC is +5v. Atmega1280 has 86 GPIO, 11 GPIO ports. They are PORTA, PORTB, PORTC, PORTD, PORTE, PORTF, PORTG, PORTH, PORTJ, PORTK and PORTL. Among these ports PORTA, PORTB, PORTF, PORTJ, PORTK, PORTL, PORTH and PORTG are all input. PORTC, PORTD, PORTE & PORTG5_bit are all output. To give input toggle switches are used and to see the output LEDs are used. There are two more inputs one is Enable and another is for clock. Microcontroller has its own clock connected to Xtal1 & Xtal2 pin. Value of oscillator is 16 MHz. A reset circuit is connected with RESET_bar pin to reset the microcontroller. This microcontroller can execute 28 operations and lots of sub operations. To select main operation PORTH0 to PORTH5 total 6 pin is used and to select sub operation PORTG0 to PORTG4 is used. To Enable and disable the IC, Enable_bar is used at PORTH6_bit. External clock is used at PORTH7_bit. This IC can operate in both asynchronous and synchronous mode. When there is no external clock it will operate in asynchronous mode and while clock is used it will operate based on the clock frequency. If Enable_bar pin is high at that time

microcontroller will not take any clock input or will not work in asynchronous mode also.

a) Main Technology Used

Main technology of this IC is the programming method. To perform any operation with this IC, at first the main operation mode and sub operation mode is selected. For example if we want to do an addition operation the mode 00 is selected. PortH0 to PortH5 has to be 0. Then we have to select sub operation like we want to do an 8bit, 16 bit or 24bit addition. If we want an 8bit addition we have to set the value of PortG0_bit & PORTG1_bit equal to 0. As this is an 8bit operation we have the option to choose with how many variable we would like to do the operation. We can choose an addition operation from two variables to six variables maximum. Let's choose 2 variables. Then Value of PortG2_bit is 0, PortG3_bit is 1 & PortG3_bit is 0. Here PortG2_bit is LSB and PortG3_bit is MSB. Now it is ready to perform two 8 bit additions. Any value of PortA will be added with any value of PortB and output will be shown by LEDs in binary format. PortC0_bit is the LSB bit of output and PortG5_bit is the MSB bit.

For an 8bit operation we need to select number of inputs but for 16bit or 24bit operation this IC will start working instantly when main operation and sub

operation is set. There are some operations in 16 bit or 24 bit which have sub operations also. For example if we want to execute a 24bit Barrel shifter, we have to set the main operation first. H5=0, H4=1, H3=1, H2=0, H1=1, H0=0. Then we have to select sub operation that, shift it right or shift it left. So, the sub operation G0=0, G1=0 is set for shifting right arithmetic. Then we have to select the amount of shifting through Pin G2 to G4. The table below gives us the operation chart and bit values to set them

Table 1 : Main Operation Chart

	H4	H3	H2	H1	H0	Operation
00	0	0	0	0	0	Addition
01	0	0	0	0	1	Subtraction
02	0	0	0	1	0	Multiplication
03	0	0	0	1	1	Division
04	0	0	1	0	0	AND
05	0	0	1	0	1	OR
06	0	0	1	1	0	NOT
07	0	0	1	1	1	NOR
08	0	1	0	0	0	NAND
09	0	1	0	0	1	XOR
10	0	1	0	1	0	XNOR
11	0	1	0	1	1	Buffer
12	0	1	1	0	0	Universal Shift Register
13	0	1	1	0	1	Ones Counter
14	0	1	1	1	0	Binary Counter
15	0	1	1	1	1	Inverted Counter
16	1	0	0	0	0	Mux

17	1	0	0	0	1	Demux
18	1	0	0	1	0	Segment Decoder Common Cathode
19	1	0	0	1	1	Segment Decoder CC Dot
20	1	0	1	0	0	Segment Decoder Common Anode
21	1	0	1	0	1	Seg Dec CA Dot
22	1	0	1	1	0	Priority Encoder
23	1	0	1	1	1	Priority Decoder
24	1	1	0	0	0	Parity Checker
25	1	1	0	0	1	Comparator
26	1	1	0	1	0	Barrel Shifter
27	1	1	0	1	1	Ring Counter
28	1	1	1	0	0	John. Ring Counter

Table 2 : Sub operation chart

G1	G0	Sub Operation	Main Operation
0	0	8 bit	Addition, Subtraction, Multiplication, Division, All operations
0	1	16 bit	
1	0	24 bit	

Table 3 : Sub operation chart

G1	G0	Sub Operation	Main Operation
0	0	7 Segment	Segment Display
0	1	14 Segment	
1	0	16 Segment	

Table 4 : 8 bit level operations

G4	G3	G2	IN1	IN2	IN3	IN4	IN5	IN6	Equation	Operation
0	1	0	PortA	PortB	X	X	X	X	IN1+IN2	ADD
0	1	1	PortA	PortB	PortF	X	X	X	IN1+IN2+IN3	
1	0	0	PortA	PortB	PortF	PortJ	X	X	IN1+IN2+IN3+IN4	
1	0	1	PortA	PortB	PortF	PortJ	PortK	X	IN1+IN2+IN3+IN4+IN5	
1	1	0	PortA	PortB	PortF	PortJ	PortK	PortL	IN1+IN2+IN3+IN4+IN5+IN6	

Table 5 : 16 bit level operations

G4	G3	G2	IN1	IN2	IN3	Equation	Operation
0	1	0	PortB PortA	PortJ PortF	X	IN1+IN2	ADD
0	1	1	PortB PortA	PortJ PortF	PortL PortK	IN1+IN2+IN3	

Table 6 : 24 bit level operations

G4	G3	G2	IN1	IN2	Equation	Operation
x	x	x	PortF PortB PortA	PortL PortK PortJ	IN1+IN2	ADD

In the above table sub operation of addition is shown and how to select the operation bit values are also shown. In this way all the other operations are performed. If we want to do subtraction IN1 has to be the largest number than second largest should be IN2, so that we don't get any wrong answer for subtraction. This order is maintained for division also. For two inputs,

the equation is IN1/ IN2. Segment Decoder takes input as binary value and gives output for common cathode without dot, common cathode with dot, common anode without dot and common anode with dot. For example if L3=0, L2=1, L1=0, L0=0 then segment with show 4 as output. Ones counter, binary counter, Mux, Demux all these operations are also divided according to the

above table. But Universal shifter and barrel shifter have other sub functions. These functions are shown below:

Table 7 : Sub operation chart

L7	L6	L5	SUB	MAIN
0	0	X	Serial In Parallel Out	Universal Shift Register
0	1	X	Serial In Serial Out	
1	0	X	Parallel In Serial Out	
1	1	X	Parallel In Parallel Out	
0	0	0	Shift Right Logic	Barrel Shifter
0	0	1	Shift Right Arithmetic	
0	1	X	Rotate Right	
1	0	0	Shift Left Logic	
1	0	1	Shift Left Arithmetic	
1	1	X	Rotate Left	

Table 8 : Commands of Sub operation chart

L3	L2	L1	L0	OPERATIONS	SUB
X	0	0	0	Hold	SIPO SISO PISO PIPO
X	0	0	1	Load	
X	0	1	0	Load & Shift Right	
X	0	1	1	Load & Shift Left	
X	1	0	0	Shift Circular Right	
X	1	0	1	Shift Circular Left	
X	1	1	0	Shift Arith Right	
X	1	1	1	Shift Arith Left	
0	0	0	0	Shift 0	Barrel Shifter
0	0	0	1	Shift 1	
0	0	1	0	Shift 2	
0	0	1	1	Shift 3	
0	1	0	0	Shift 4	
0	1	0	1	Shift 5	
0	1	1	0	Shift 6	
0	1	1	1	Shift 7	
...	...	...	...		
1	1	1	1	Shift 15	

b) Device Used

Brain of this project is Atmega1280 microcontroller. It is an 8 bit Micro controller with RISC architecture. Its speed is up to 16 MIPS throughout at 16MHz. It has 128K bytes of flash and 4Kbytes EEPROM. Operating voltage 2.7v -5.5v, in active mode it consumes only 1.1mA & in sleep mode it consumes less than 1uA current. It has 86 GPIO, 16 PWM channels, SPI, 4 UART, I2C and other features which made it a perfect choice of designing a versatile digital design IC.

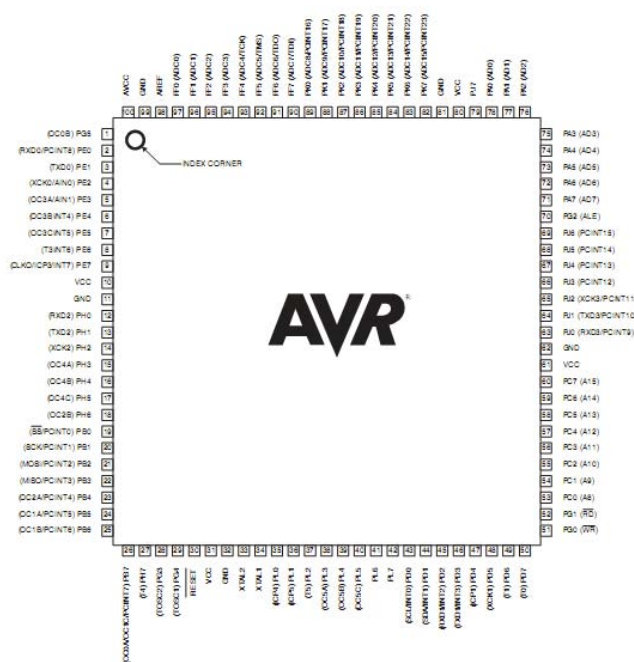


Figure 3 : Atmega1280 Microcontroller[8]

Our program consumes only 37kbytes of Flash memory. Still 91kbytes of Flash memory is available to design other operations.

As the outputs are shown in simulation no resistors has been used. In practical implementation with each LED 220ohm resistance is being used. Each input Pin has been pulled down through 10k resistor.

Atmega1280 contains 135 powerful instructions like addition, subtraction, shifting etc and most single clock cycle execution. 32x8 general purpose working registers, two 8bit timer & four 16bit timer, interrupt and wakeup on pin change. It comes in 100 lead TQFP, 100 Ball CBGA which can be breakout easily and small in size. This microcontroller has the most number of GPIO and all 86 GPIO pins can be used as input or output. Brown out feature of this IC makes it more stable output and long lasting.

III. TEST RESULTS

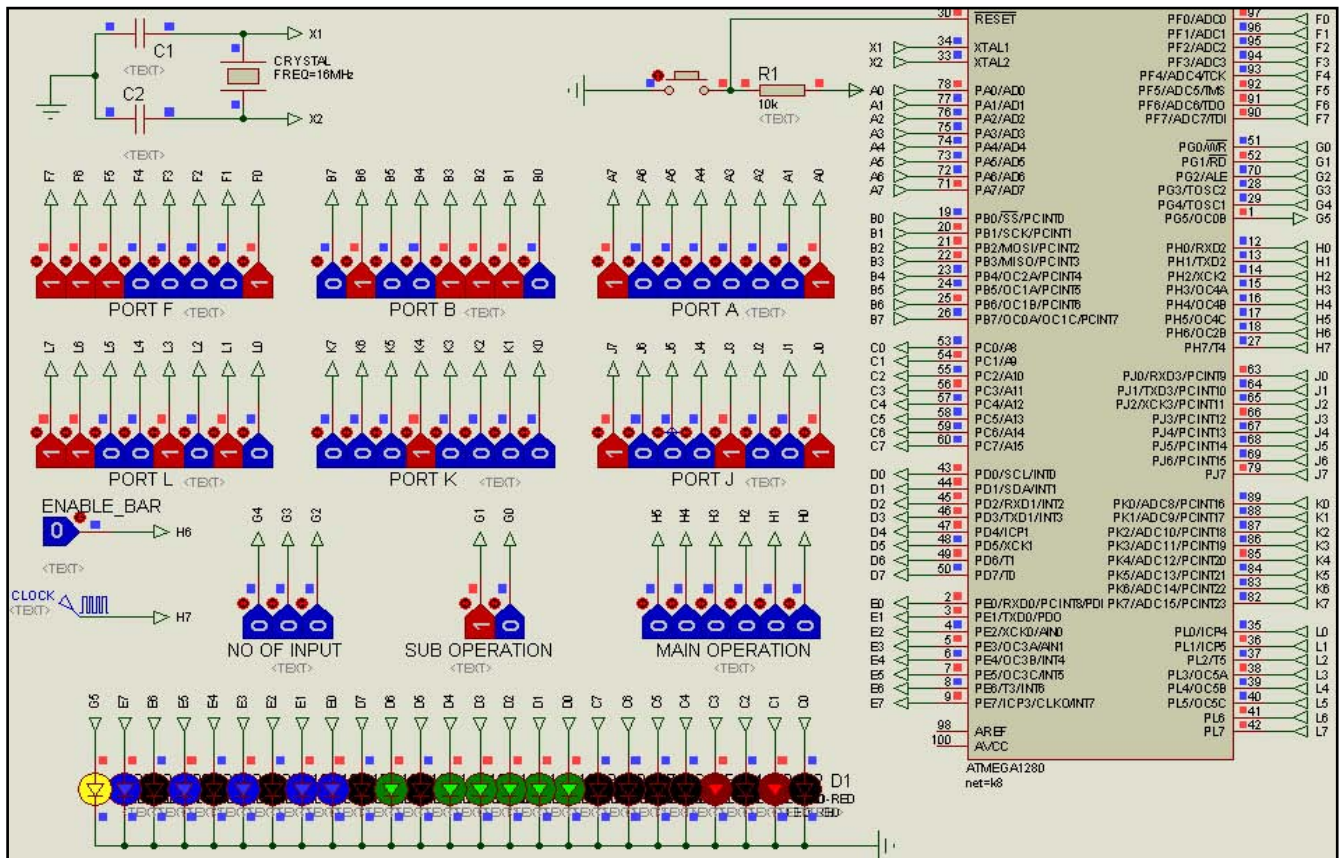


Figure 4 : 24bit Addition Operation[10]

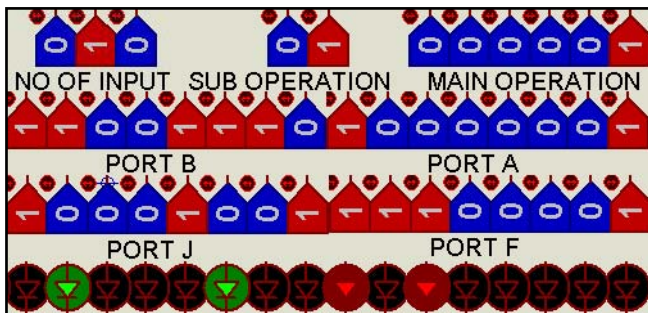


Figure 5 : 16bit 2input Subtraction Operation

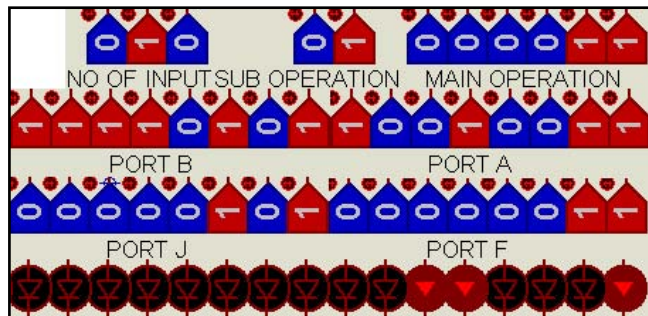


Figure 6 : 16bit 2input Division Operation

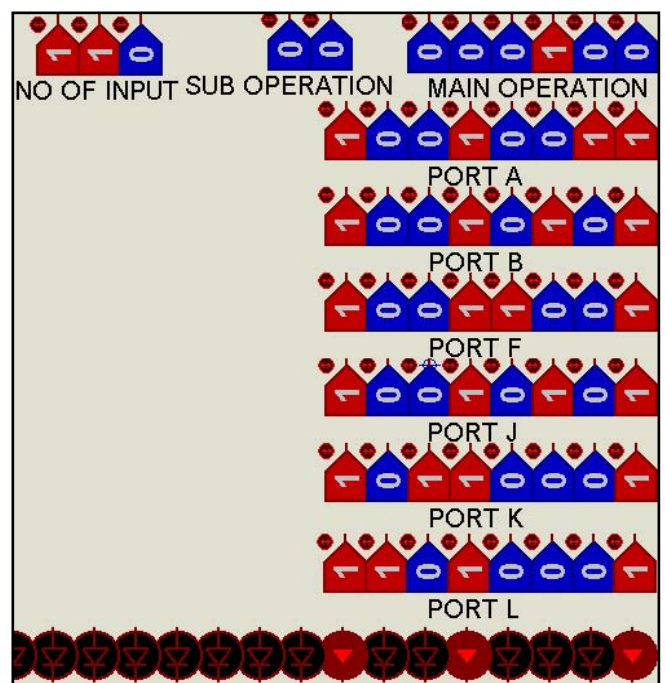


Figure 7 : 8bit 6input AND Operation

IV. FURTHER APPLICATIONS

- To practice large truth table up to 6 variables this IC can become handy.
- This IC can be used in many complex Digital design systems.
- Students can learn different type of Digital logic operation through only one IC.
- This IC can be used where expensive digital logic design lab can't be set up for education purpose.
- Rapid prototype for teaching Digital Logic Design.

REFERENCES RÉFÉRENCES REFERENCIAS

1. Md. Khaled Hossain and Md. Nasimuzzaman Chowdhury, International Journal of Emerging Technology and Advanced Engineering, Volume 3, Issue 6, June 2013, PP. 34-41.
2. Stanisavljevic, Z.; Pavlovic, V.; Nikolic, B.; Djordjevic, J., "SDLDS- System for Digital Logic Design & Simulation", Education, IEEE Transactions on (Volume: 56 , Issue: 2), May 2013.
3. Zemva, A.; Trost, A.; Zajc, B. "A Rapid Prototyping Environment for teaching Digital Logic Design", Education, IEEE Transactions on (Volume: 41, Issue: 4), Nov 1998.
4. Digital design principles and practices by John F. Wakerly.
5. http://www.electronicstutorials.ws/sequential/seq_5.html for details operation of Shift Register.
6. <http://tams-www.informatik.uni-hamburg.de/applets/hades/webdemos/20-arithmetic/10-adders/halfadd-fulladd.html> for lots of digital design simulation verification.
7. Guillermo A. Vera, Jorge Parra, Craig Kief, Marios Pattichis and Howard Pollard, "Integrating Reconfigurable Logic in the First Digital Logic Course", 9th International Conference on Engineering Education, July 23-28 2006.
8. Poulastya Mukherjee. Remote Control of Industrial Machines Using Mobile Phone. IJETAE, Volume 3, Issue 1, January 2013 for project outline diagram.
9. Atmega1280 datasheet by Atmel.
10. www.mikroe.com for free demo compiler and help of compiler software.
11. http://www.labcenter.com/products/vsm/vsm_overview.cfm for simulator and demo videos of simulation.

GLOBAL JOURNALS INC. (US) GUIDELINES HANDBOOK 2013

WWW.GLOBALJOURNALS.ORG

FELLOW OF ASSOCIATION OF RESEARCH SOCIETY IN COMPUTING (FARSC)

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

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

MEMBER OF ASSOCIATION OF RESEARCH SOCIETY IN COMPUTING (MARSC)

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

AUXILIARY MEMBERSHIPS

ANNUAL MEMBER

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

PAPER PUBLICATION

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



PROCESS OF SUBMISSION OF RESEARCH PAPER

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

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

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

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

(II) Choose corresponding Journal.

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

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

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

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



PREFERRED AUTHOR GUIDELINES

MANUSCRIPT STYLE INSTRUCTION (Must be strictly followed)

Page Size: 8.27" X 11"

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

You can use your own standard format also.

Author Guidelines:

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

1. GENERAL

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

Scope

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

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

2. ETHICAL GUIDELINES

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

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

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

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

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

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

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

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

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

Please mention proper reference and appropriate acknowledgements wherever expected.

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

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

3. SUBMISSION OF MANUSCRIPTS

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

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



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

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

4. MANUSCRIPT'S CATEGORY

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

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

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

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

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

5. STRUCTURE AND FORMAT OF MANUSCRIPT

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

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

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

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



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

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

Format

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

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

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

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

Structure

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

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

Abstract, used in Original Papers and Reviews:

Optimizing Abstract for Search Engines

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

Key Words

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

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

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

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



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

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

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

Acknowledgements: Please make these as concise as possible.

References

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

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

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

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

Tables, Figures and Figure Legends

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

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

Preparation of Electronic Figures for Publication

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

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

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



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

6. AFTER ACCEPTANCE

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

6.1 Proof Corrections

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

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

(Free of charge) from the following website:

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

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

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

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

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

6.3 Author Services

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

6.4 Author Material Archive Policy

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

6.5 Offprint and Extra Copies

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

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

Introduction:

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

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

Approach:

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



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

Procedures (Methods and Materials):

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

Materials:

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

Methods:

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

Approach:

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

What to keep away from

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

Results:

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

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



Content

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

What to stay away from

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

Approach

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

Figures and tables

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

Discussion:

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

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

Approach:

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



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

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

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

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



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

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

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



INDEX

A

Adenine · 40
Amenable · 36
Anisotropy · 21
Attenuated · 25

C

Canonical · 18
Convolution · 1

D

Diligent · 23

F

Fidelity · 24

G

Guanine · 40

M

Mitigated · 35

N

Nucleotides · 40

O

Overwhelming · 36

P

Protuberant · 16

S

Spurious · 25
Symposium · 14, 47

T

Trellis · 1, 4, 5, 6

V

Verilog · 32
Viterbi · 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 15

X

Xilinx · 1, 11, 13, 23, 27, 32, 33

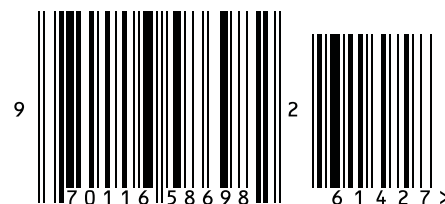


save our planet



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

© Global Journals Inc.