

Reducing Shared Cache Misses via Dynamic Coscheduling on Multicores

by

Wael Amr Hossam El Din Lotfy

A Thesis Submitted to the
Faculty of Engineering at Cairo University
In Partial Fulfillment of the Requirements for the Degree of

MASTER OF SCIENCE
in
COMPUTER ENGINEERING

FACULTY OF ENGINEERING, CAIRO UNIVERSITY
GIZA, EGYPT
2015

Reducing Shared Cache Misses via Dynamic Coscheduling on Multicores

by

Wael Amr Hossam El Din Lotfy

A Thesis Submitted to the
Faculty of Engineering at Cairo University
In Partial Fulfillment of the Requirements for the Degree of

MASTER OF SCIENCE
in
COMPUTER ENGINEERING

Under the Supervision of

Dr. Ihab ElSayed Talkhan
Professor
Computer Engineering Department
Cairo University

Dr. Hany Mohamed ElSayed
Assistant Professor
Communication Engineering Department
Cairo University

FACULTY OF ENGINEERING, CAIRO UNIVERSITY
GIZA, EGYPT
2015

Reducing Shared Cache Misses via Dynamic Coscheduling on Multicores

by

Wael Amr Hossam El Din Lotfy

A Thesis Submitted to the

Faculty of Engineering at Cairo University

In Partial Fulfillment of the Requirements for the Degree of

MASTER OF SCIENCE
in
COMPUTER ENGINEERING

Approved by the examining committee:

Prof. Dr. Ihab ElSayed Talkhan,

Thesis Main Advisor

Associate Prof. Dr. Ibrahim Mohamed Kamar,

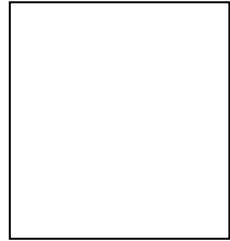
Internal Examiner

Prof. Dr. Ahmed Abdelrahman Abou-Auf,

External Examiner

FACULTY OF ENGINEERING, CAIRO UNIVERSITY
GIZA, EGYPT
2015

Engineer: Wael Amr Hossam El Din Lotfy
Date of Birth: 07 / 07 / 1984
Nationality: Egyptian
E-mail: wael.amrhossam@gmail.com
Phone: 01285692539
Address: Haydiek El Ahram, Masakin El Sharkaa El Watania, El Haram.
Registration Date: 01 / 10 / 2008
Awarding Date: / /
Degree: Masters of Science
Department: Computer Engineering



Supervisors: Prof. Dr. Ihab ElSayed Talkhan
Assistant Prof. Dr. Hany Mohamed ElSayed

Examiners: Prof. Dr. Ihab E. Talkhan (Cairo University)
Associate Prof. Dr. Ibrahim M. Qamar (Cairo University)
Prof. Dr. Ahmed A. Abou-Auf (American University)

Title of Thesis:
Reducing Shared Cache Misses Via Dynamic Coscheduling On Multicores

Keywords:
Cache Miss Rate, Dynamic Scheduling, Multicore, Symbiosis

Summary:

Multicore technology enables having two processors or more working together on the same single chip. It enables the system to perform more tasks with higher overall system performance. However, this performance can't be exploited well due to the high miss rate in the second level shared cache among the cores which represents one of the multicore's challenges. This problem can be reduced by using the scheduling techniques. This thesis addresses the dynamic co-scheduling of tasks in multicore real-time systems. The focus is on the basic idea of the megatask technique for grouping the tasks that may affect the shared cache miss rate, and the Pfair scheduling that is then used for reducing the concurrency within the grouped tasks while ensuring the real time constraints. So all the grouped tasks could not be run at the same time. Consequently the shared cache miss rate is reduced. The dynamic co-scheduling is proposed through the combination of the symbiotic technique with the megatask technique for co-scheduling the tasks based on the collected information.

Experiments show that the proposed dynamic co-scheduling can decrease the shared cache miss rate compared to the static one by up to 52%.

Advisor's Signature:

Acknowledgments

First of all, I would like to thank ALLAH for his guidance and help. While I wrote this thesis, it is ultimately the result of inspiration from, and interaction with, a large number of people. I thank my advisors, Dr. Hany Mohamed Elsayed and Dr. Ihab Elsayed Talkhan, for their support, guidance, and patience throughout my graduate career.

I also thank Islam Mohamed Hatem Atta for his supporting and providing me valuable feedback on my research.

Finally, I would like to thank my family and friends for being supportive in all that I do and for their prayers; particularly my father Amr HossamElDin, my mother, Eman Salem, my sister ,Ghadah Amr and my friend, Wael Nabil.

Dedication

I dedicate this thesis to every student who cares about the research and the knowledge for developing better nations and achieving a better humanitarian level.

Table of Contents

ACKNOWLEDGMENTS	V
DEDICATION.....	VI
TABLE OF CONTENTS.....	VII
LIST OF TABLES	IX
LIST OF FIGURES	X
ABSTRACT.....	XII
1 INTRODUCTION	1
1.1 CHIP MULTIPROCESSORS	1
1.2 THESIS PROBLEM.....	6
1.3 THESIS ASSUMPTIONS.....	6
2 LITERATURE REVIEW	7
2.1 RELATED WORK	7
2.2 CONCLUSION	12
3 GROUPING AND SCHEDULING TASKS ON MULTICORE PLATFORM	13
3.1 GROUPING:MEGATASKS TECHNIQUE.....	14
3.1.1 Packing in Megatasks	14
3.1.1.1 Concept	14
3.1.1.2 Cumulative Weight	15
3.1.1.3 Packing Strategy	15
3.1.2 Re-weighting Rules	16
3.1.2.1 Deadlines.....	16
3.1.2.2 Correctness Proof.....	16
3.1.2.3 Re-weighting a Megatask	17
3.1.2.4 Re-weighting Example.....	18
3.2 PFAIR SCHEDULING	19
3.2.1 Lag Constraints.....	20
3.2.2 Tasks' Division	21
3.2.3 Example	21
3.2.4 Fairness.....	22
3.2.5 Optimality.....	22
3.2.6 Tie-breaking rules.....	23
3.2.6.1 First tie-break:The successor bit $b(T)$	23
3.2.6.2 Second tie break:The group deadline $D(T)$	24
3.2.7 The PD2 Priority Definition	25
3.3 SYMBIOSIS FACTOR	25
3.3.1 Concept.....	26
3.3.2 Sampling.....	26
3.3.3 Probabilistic Modeling	27
3.3.4 Applying Symbiosis to our research topic “Multicore platforms”	27

3.3.4.1 Classification Scheme	28
3.3.4.1.1 An “Animalistic” Taxonomy	29
3.3.4.1.2 SDC	30
3.3.4.1.3 Pain.....	30
3.3.4.1.4 Miss Rate.....	31
3.3.4.1.5 The Proposed Classification Scheme “Temporal Working Set Size”	31
3.4 DYNAMIC GROUPING	32
3.5 CONTRIBUTION	33
4 SIMULATION METHODOLOGY	34
4.1 SESC SIMULATOR	34
4.1.1 Evaluation.....	34
4.2 GEM5 SIMULATOR.....	34
4.2.1 Evaluation.....	35
4.3 CACHE SIMULATOR	35
4.3.1 Design Phases.....	35
4.3.1.1 First phase ”Memory Trace Collection”:.....	36
4.3.1.2 Second phase ”Build Our Proposed Scheduler”:	36
4.3.1.2.1 First Configuration ”Pfair without grouping”	36
4.3.1.2.2 Second Configuration “Static Megatask based on Working Set Size”	36
4.3.1.2.3 Third Configuration “Dynamic Megatask based on Working Set Size”	37
4.3.1.2.4 Fourth Configuration“Dynamic Megatask based on Miss Rate”.....	37
4.3.1.3 Third Phase “Cache Simulator”:.....	37
4.3.1.3.1 Operational Scenario	38
4.3.1.4 Fourth Phase “Test Cases”:.....	39
5 SIMULATION RESULTS	40
5.1 TESTING SCENARIOS.....	40
5.1.1 Scenario 1	41
5.1.2 Scenario 2	43
5.1.3 Scenario 3	44
5.1.4 Scenario 4	46
5.1.5 Scenario 5	48
5.1.6 Scenario 6	50
5.1.7 Scenario 7	52
5.1.8 Observation.....	54
6 CONCLUSION AND FUTURE WORK	55
6.1 CONCLUSION	55
6.2 FUTURE WORK	56

List of Tables

Table 5.1 L1 Cache and Main Memory Parameters	40
Table 5.2 Scenario 1 Parameters.....	41
Table 5.3 Scenario 2 Parameters.....	43
Table 5.4 Scenario 3 Parameters.....	44
Table 5.5 Scenario 4 Parameters	46
Table 5.6 Scenario 5 Parameters.....	48
Table 5.7 Scenario 6 Parameters.....	50
Table 5.8 Scenario 7 Parameters.....	52
Table 5.9 Improved Shared Cache L2 Miss Rate	54

List of Figures

Figure 1.1 Conventional Microprocessor	2
Figure 1.2 Simple Chip Multiprocessor.....	2
Figure 1.3 Shared Cache Multiprocessor.....	3
Figure 1.4 Multithreaded Shared Cache Multiprocessor	3
Figure 1.5 Multi Chip Module	4
Figure 3.1 The Proposed Dynamic Co-scheduling Technique	13
Figure 3.2 Packing in Megatasks	14
Figure 3.3 Re-weighting Rules	16
Figure 3.4 Run Pfair Scheduler.....	19
Figure 3.5 Three Tasks Sets.....	21
Figure 3.6 Windows of the first 16 subtasks of Task T	23
Figure 3.7 Symbiosis Factor	25
Figure 3.8 Classification Scheme	28
Figure 3.9 Re-packing in Megatasks	32
Figure 4.1 Cache Simulator	35
Figure 4.2 Operational Scenario	38
Figure 5.1 : 8 SPECjvm2008 benchmarks of total size 16 MB that share L2 cache of size 1 MB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate.	42
Figure 5.2 : 4 SPECjvm2008 benchmarks of total size 8 MB that share L2 cache of size 8 MB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate.	44
Figure 5.3 : 5 SPECjvm2008 benchmarks of total size 10 MB that share L2 cache of size 20 MB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate.	45
Figure 5.4 : 5 SPECjvm2008 benchmarks of total size 10 MB that share L2 cache of size 1 MB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate.	47
Figure 5.5 : 16 SPECjvm2008 benchmarks of total size 32 MB that share L2 cache of size 1 MB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate.	49

Figure 5.6 : 12 SPECjvm2008 benchmarks of total size 24 MB that share L2 cache of size 512 KB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate. 51

Figure 5.7: 10 SPECjvm2008 benchmarks of total size 20 MB that share L2 cache of size 512 KB, X-axis represents the fetched instructions and Y-axis represents the shared cache L2 Miss Rate. 53

Abstract

Multicore technology enables having two processors or more working together on the same single chip. It enables the system to perform more tasks with higher overall system performance. However, this performance can't be exploited well due to the high miss rate in the second level shared cache among the cores which represents one of the multicore's challenges. This problem can be reduced by using the scheduling techniques. Although the scheduling techniques have been recently proposed for multicores platforms, most of them have aimed at the system throughput or cache performance without ensuring the real time constraints or measuring the shared cache miss rate. A recent work targeted scheduling for reducing cache miss rate while ensuring the real time constraints among Multicore platforms. However, this scheduling was static, i.e. it does not adapt itself to operating conditions.

This thesis addresses the dynamic co-scheduling of tasks in multicore real-time systems. The focus is on the basic idea of the megatask technique for grouping the tasks that may affect the shared cache miss rate, and the Pfair scheduling that is then used for reducing the concurrency within the grouped tasks while ensuring the real time constraints. So all the grouped tasks could not be run at the same time. Consequently the shared cache miss rate is reduced. The dynamic co-scheduling is proposed through the combination of the symbiotic technique with the megatask technique for co-scheduling the tasks based on the collected information using two schemes. The first scheme is measuring the temporal working set size of each running task at run time, while the second scheme is collecting the shared cache miss rate of each running task at run time.

The used simulation methodology is based on implementing a cache simulator that is based on SESC simulator. The PIN tool is used to generate memory access request files using different applications from the SPECJVM2008 benchmarks, then these files are used as an input for the cache simulator.

Experiments show that the proposed dynamic co-scheduling can decrease the shared cache miss rate compared to the static one by up to 52%. This indicates that the dynamic co-scheduling is important to achieve high performance with shared cache memory for running high workloads like multimedia applications that require real-time response and continuous-media data types.

1 Introduction

In 1965, Gordon E. Moore, the intel co-founder, explained the growth of semiconductors technology in a formula. This formula states that the number of transistors on a chip will double approximately every two years. This observation has since been dubbed "Moore's Law" and is now enormously influential. "Moore's Law" has been considered in the processor industry. As transistors are increased on a chip, the processor clock frequencies are also increased at the same time. "Moore's law" has reached its limit in 2010. One reason behind this is that the faster processors have power consumption and thermal dissipation. This needs special cooling systems and leads to higher cost. Hence processor industry has moved towards the multicore technology since the delivered performance of single cores can not meet the needed requirements for running different applications like web servers, multimedia programs and databases. Multicore technology is introduced to increase the required performance and power efficiency. However, there are new challenges for this technology.

Next section contains some figures that show the architectures for the uniprocessor, multiprocessors and multicores, then it tackles the challenges of the multicores. These challenges open the door in research to propose new software and hardware solutions.

1.1 Chip Multiprocessors

Chip Multiprocessors (can also be named a Multicore processor) refers to a single chip that integrates two or more processors in an area that would have originally been filled with a single large uniprocessor. This solves the power consumption problem when adding more transistors to the uniprocessor and switching it at higher and higher frequencies.

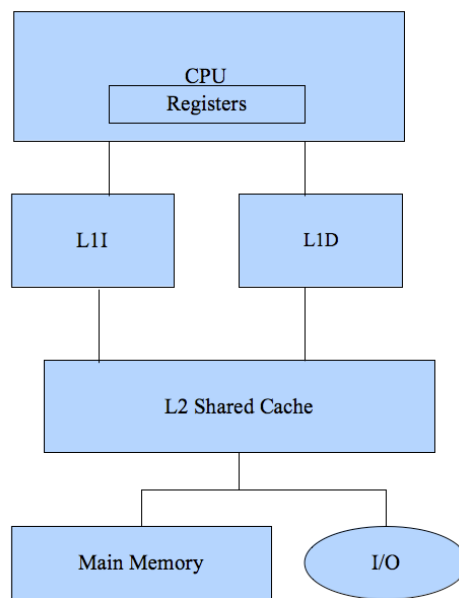


Figure 1.1 Conventional Microprocessor

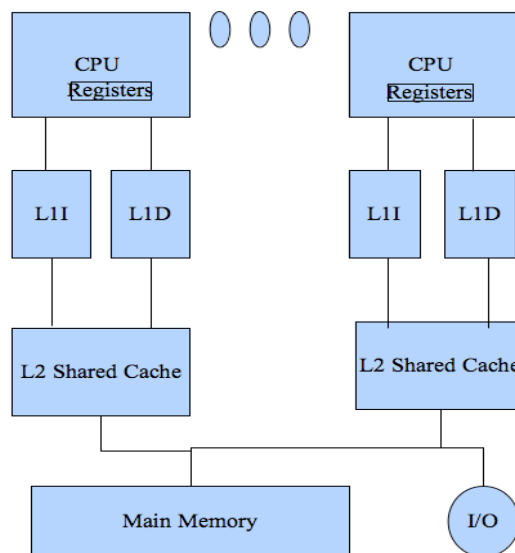


Figure 1.2 Simple Chip Multiprocessor

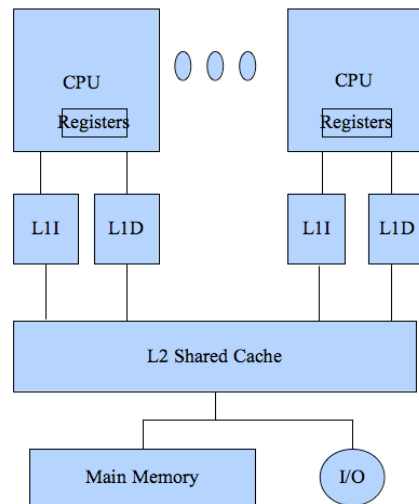


Figure 1.3 Shared Cache Multiprocessor

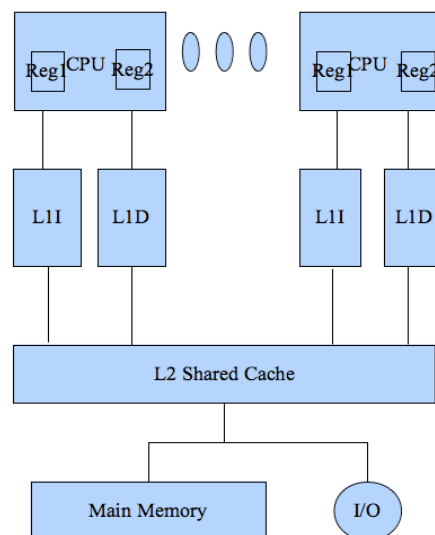


Figure 1.4 Multithreaded Shared Cache Multiprocessor