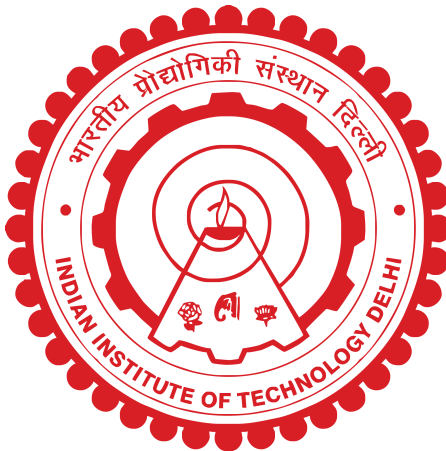


LEAKAGE AWARE DYNAMIC THERMAL MANAGEMENT FOR 3D MEMORY ARCHITECTURES

LOKESH SIDDHU



**DEPARTMENT OF COMPUTER SCIENCE & ENGINEERING
INDIAN INSTITUTE OF TECHNOLOGY DELHI**

December 2022

© **Indian Institute of Technology Delhi (IITD), New Delhi, 2022**

LEAKAGE AWARE DYNAMIC THERMAL MANAGEMENT FOR 3D MEMORY ARCHITECTURES

by

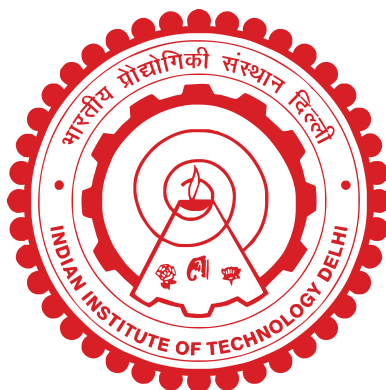
Lokesh Siddhu

Department of Computer Science and Engineering

Submitted

in fulfillment of the requirements of the degree of Doctor of Philosophy

to the



INDIAN INSTITUTE OF TECHNOLOGY DELHI

December 2022

Certificate

This is to certify that the thesis titled “**Leakage Aware Dynamic Thermal Management for 3D Memory Architectures**” being submitted by **Lokesh Siddhu** for the award of **Doctor of Philosophy** in Computer Science and Engineering is a record of bonafide work carried out by him under my guidance and supervision in the Department of Computer Science and Engineering, Indian Institute of Technology Delhi. The work presented in this thesis has not been submitted elsewhere, either in part or full, for the award of any other degree or diploma unless otherwise stated explicitly.

Certain works included in this thesis involved collaboration or use of measurement data obtained by other students, which have been explicitly specified/acknowledged in the corresponding chapters and the part done by those collaborators appeared in their respective thesis.

Preeti Ranjan Panda

Professor

Department of Computer Science and Engg.

Indian Institute of Technology Delhi

New Delhi- 110016

Acknowledgments

“Gratitude is not an attitude. Gratitude is something that flows out of you when you are overwhelmed by the recognition of what you have received.”

—Sadhguru

I would like to express my gratitude to Prof. Preeti Ranjan Panda for his support, guidance, and insight. His unwavering commitment to quality has deeply touched me. Also, his keen attention, careful listening, and understanding are qualities worth emulating in every aspect of life. His constructive feedback on various presentation aspects has improved my skills by leaps and bounds. Prof. Panda has been a joy to share time and work with.

I have had the privilege of getting insightful and valuable feedback from Prof. Anshul Kumar, Prof. Kolin Paul, and Prof. Mukul Sarkar throughout my research. Their suggestions, understanding, and critique have been very helpful. Prof. M. Balakrishnan’s vision has been a constant source of inspiration. His ability to present complex problems clearly and simply provided ease in understanding and imbibing. Also, I had a wonderful time TA-ing. It offered me the experience of conducting and managing labs and marking examinations and assignments. I would like to thank all the related professors.

Rajesh Kedia has been a wonderful friend, collaborator, and colleague at IIT Delhi. Almost daily, we discussed research problems and solutions, resulting in a regular knowledge update. His support, both in personal and professional capacities, has been of immense value to me. Over the years, I have had the good fortune of collaborating with many professors and students - Anuj Pathania, Rajesh Kedia, Shailja Pandey, Aritra Bagchi, Martin Rapp, Isaar Ahmad, Hameedah Sultan. Their ideas and suggestions have helped me thoroughly. Discussions with them have resulted in improved research quality and presentation. I would also like to thank the members of the Memory and Architecture Research Group (MARG) - Rajesh Kedia, Neetu Jindal, Shailja Pandey, Sakshi Tiwari, Hadi Brais, Divya Praneetha, Isaar Ahmad, Aritra Bagchi, Ayushi Agarwal, Garima Modi, Zilmarij Iqbal, for their everlasting support and company. The technical feedback in our weekly meeting ensured steady progress and learning. In addition, our informal conversations over lunch/tea outings provided a daily dose of humor, joy, and research ideas.

In the Ph.D. journey, my labmates have played a crucial role in creating a conducive atmosphere for research. I would like to thank Kunal Dahiya, Ankit Anand, Dinesh Khandelwal, Vishal Sharma, Dilpreet Kaur, Harsh Das, Samuel Wedaj, Solomon Abera, Richa Gupta, Himanshu Gandhi and Iqra Altaf. Our discussion ranged from technical topics to

global leadership and polity, always resulting in a more thorough understanding of research or any other situation. Many parts of my research's foundation were established with the help of Sandeep Chandran, Namita Sharma, Rajshekar Kalayappan, and Prathmesh Kallurkar. Time spent with them was invaluable.

During my Ph.D., I had the opportunity to organize volunteer activities for VLSID 2019. It was an immensely enriching experience. In many ways, it depicts how the crux of the academic culture lies in several volunteers. Various selfless volunteers have built the community, which allows us to function. I wish to thank them from the bottom of my heart.

The office staff at Amar Nath and Shashi Khosla School of Information Technology and the Department of Computer Science have been most helpful. I would like to thank Ms. Vandana, Ms. Rekha, Mr. Hemant, Mr. Suresh, Mr. Rajesh, Mr. Som Dutt, and many others - they have been most helpful and streamlined many administrative activities. My stay at the hostel placed me in a youthful and exuberant environment. Providing a stay and food on the campus is a luxury. I thank the hostel for their diligent effort. Much of my research assistance was funded by the Ministry of Electronics and Information Technology (MeitY) with the Visvesvaraya Ph.D. Scheme. Their support and efficiency are commendable.

I am fortunate to have friends connected dearly. Every moment spent with them was a joy. Immense gratitude to Ankit Gupta, Swati Gupta, Gaurav Singh, Disha Singh, Ruchik Shastri, Manvi Sharma, Amit Mishra, Namrata Mishra, Srijith Haridas, Nagraj Swami, Karthika Ravindran, Sumeet Mali, Nimash Choudhary, Chinmay Jain, Mehul Agrawal. Their feedback and understanding have positively shaped various aspects of my thought and emotion.

My family has been a constant source of support. Their commitment to education, learning, and knowledge has been exemplary. My parents have especially done their best to take care of me and provide a comfortable environment for my pursuit of knowledge. As I reflect, life is connected in several myriad ways, and countless people have contributed to enabling me. All the masters over generations have passed their knowledge and wisdom. My sincere gratitude and love to all life. It has been an immensely fulfilling and joyful journey.

Lokesh Siddhu

Abstract

Memory architects have relied on process scaling (making transistors, wires smaller) to gain transistor density and performance in the last few decades. However, the growing fabrication/production costs and variability in the manufacturing process have motivated them to explore 3D architectures, which allow the stacking of 2D memories to gain an order of magnitude of bandwidth. However, the increased stacking and power density in 3D memories make heat dissipation a challenge. At elevated temperatures, the transistors (and chips) may operate unreliably. Higher temperatures than the safe operating temperature might also cause permanent damage to the chip. Cooling costs increase significantly when on-chip temperatures are higher. Leakage current grows exponentially with temperature causing a further rise in temperature and may lead to thermal runaway. Hence, there is a need to study low overhead thermal-aware system-level solutions for 3D memory architectures.

While prior works have looked at reducing dynamic/switching power through reduced memory accesses, in these memories, both leakage and dynamic power consumption are comparable. Furthermore, as the temperature rises, the leakage power increases, creating a thermal-leakage loop. We present the first study of the thermal-leakage loop resulting in a system-level DTM policy, *FastCool*, based on data migration to the 2D memory. We propose a detailed analytical model to quantify the memory delay, including queuing delays for a hybrid (3D + 2D) memory system. We suggest an improved DTM strategy, *Energy-Efficient FastCool (EEFC)*, which incorporates additional designer insights to dynamically select the channels to be closed, improving the performance and energy further. We also identify that EEFC is resilient to process variability.

To lower DTM overheads, we present a proactive DTM policy, *PredictNcool*, which uses a lightweight and accurate steady-state temperature predictor for 3D memories. *PredictNcool* is the first work exploiting various designer insights about the layout/floorplan (including different symmetries) to reduce model parameters of a machine learning-based predictor, significantly reducing the required training data and prediction time. Further, we discuss low power state based DTM approaches, decreasing energy overheads, and energy-delay-product by avoiding

data migration and 2D memory usage. We suggest a DTM policy to reduce the performance overheads by integrating low power state and partial channel turn-off based approaches.

As cores and 3D memory get heated in modern computing systems, we suggest a joint thermal management approach, *CoreMemDTM*, avoiding inefficiencies caused by independently running DTMs and minimizing performance overheads. We suggest employing DTM depending on the thermal margin since safe temperature thresholds might differ for the two subsystems. We present a stall-balanced core DVFS policy for core thermal management that enables distributed cooling, decreasing overheads. Further, we introduce *CoMeT*, the integrated Core and Memory Thermal simulation toolchain, which supports various core and 3D memory configurations. The source code of *CoMeT* has been made open for public use under the *MIT* license.

तीन-डी मेमोरी आर्किटेक्चर के लिए लीकेज अवेयर डायनामिक थर्मल मैनेजमेंट

सार:

मेमोरी आर्किटेक्चर्स ने पिछले कुछ दशकों में ट्रांजिस्टर घनत्व और प्रदर्शन हासिल करने के लिए प्रोसेस स्केलिंग (ट्रांजिस्टर, तारों को छोटा बनाना) पर भरोसा किया है। हालांकि, निर्माण प्रक्रिया में बढ़ती निर्माण/उत्पादन लागत और परिवर्तनशीलता ने उन्हें तीन-डी आर्किटेक्चर का पता लगाने के लिए प्रेरित किया है, जो बैंडविड्थ की परिमाण के क्रम को प्राप्त करने के लिए दो-डी मेमोरी के स्टैकिंग की अनुमति देता है। हालांकि, तीन-डी मेमोरी में बढ़ी हुई स्टैकिंग और पावर डेंसिटी गर्मी अपव्यय को एक चुनौती बनाती है। ऊंचे तापमान पर, ट्रांजिस्टर (और चिप्स) अविश्वसनीय रूप से काम कर सकते हैं। सुरक्षित ऑपरेटिंग तापमान से अधिक तापमान भी चिप को स्थायी नुकसान पहुंचा सकता है। जब ऑन-चिप तापमान अधिक होता है तो कूलिंग की लागत काफी बढ़ जाती है। लीकेज करंट तापमान के साथ तेजी से बढ़ता है जिससे तापमान में और वृद्धि होती है और इससे थर्मल पलायन हो सकता है। इसलिए, तीन-डी मेमोरी आर्किटेक्चर के लिए कम ओवरहेड थर्मल-जागरूक सिस्टम-स्तरीय समाधानों का अध्ययन करने की आवश्यकता है।

पिछले कार्यों ने कम मेमोरी एक्सेस के माध्यम से गतिशील/स्विचिंग पावर को कम करने पर ध्यान दिया है, इन मेमोरी में, रिसाव और गतिशील बिजली की खपत दोनों तुलनीय हैं। इसके अलावा, जैसे-जैसे तापमान बढ़ता है, रिसाव की शक्ति बढ़ती है, जिससे थर्मल-रिसाव लूप बनता है। हम थर्मल-लीकेज लूप का पहला अध्ययन प्रस्तुत करते हैं जिसके परिणामस्वरूप सिस्टम-स्तरीय डीटीएम नीति, फास्टकूल, दो-डी मेमोरी में डेटा माइग्रेशन पर आधारित है। हम एक हाइब्रिड (तीन-डी + दो-डी) मेमोरी सिस्टम के लिए क्यूइंग देरी सहित मेमोरी देरी को निर्धारित करने के लिए एक विस्तृत विश्लेषणात्मक मॉडल का प्रस्ताव करते हैं। हम एक बेहतर डीटीएम रणनीति, ऊर्जा-कुशल फास्टकूल का सुझाव देते हैं, जो प्रदर्शन और ऊर्जा को और बेहतर बनाने के लिए चैनलों की गतिशील रूप से चुनने के लिए अतिरिक्त डिज़ाइनर अंतर्दृष्टि को शामिल करती है। हम यह भी पहचानते हैं कि ऊर्जा-कुशल फास्टकूल प्रक्रिया परिवर्तनशीलता की उपस्थिति में भी अच्छा प्रदर्शन करती है।

डीटीएम ओवरहेड्स को कम करने के लिए, हम एक सक्रिय डीटीएम नीति प्रस्तुत करते हैं, प्रिडिक्टएनकूल, जो तीन-डी मेमोरी के लिए हल्के और सटीक स्थिर-अवस्था तापमान पूर्वसूचक का उपयोग करती है। प्रिडिक्टएनकूल मशीन लर्निंग-आधारित प्रिडिक्टर के मॉडल मापदंडों को कम करने के लिए लेआउट/फ्लोरप्लान (विभिन्न समरूपता सहित) के बारे में विभिन्न डिज़ाइनर अंतर्दृष्टि का उपयोग करने वाला पहला काम है, आवश्यक प्रशिक्षण डेटा और भविष्यवाणी समय को महत्वपूर्ण रूप से कम करता है। इसके अलावा, हम डेटा माइग्रेशन और दो-डी मेमोरी उपयोग से बचने के द्वारा कम बिजली की स्थिति आधारित डीटीएम दृष्टिकोण, कम ऊर्जा खर्च और ऊर्जा-देरी-उत्पाद पर चर्चा करते हैं। हम कम बिजली की स्थिति और आंशिक चैनल टर्न-ऑफ आधारित दृष्टिकोणों को एकीकृत करके प्रदर्शन ओवरहेड्स को कम करने के लिए एक डीटीएम नीति का सुझाव देते हैं।

चूंकि आधुनिक कंप्यूटिंग सिस्टम में कोर और तीन-डी मेमोरी गर्म हो जाती है, हम एक संयुक्त थर्मल प्रबंधन दृष्टिकोण, कोर मेम डीटीएम का सुझाव देते हैं, ताकि स्वतंत्र रूप से डीटीएम चलाने और प्रदर्शन ओवरहेड्स को कम करने के कारण होने वाली अक्षमताओं से बचा जा सके। हम थर्मल मार्जिन के आधार पर डीटीएम को नियोजित करने का सुझाव देते हैं क्योंकि सुरक्षित तापमान थ्रेसहोल्ड दो सबसिस्टम के लिए भिन्न हो सकते हैं। हम कोर थर्मल प्रबंधन के लिए एक स्ताल-संतुलित कोर डीवीएफएस नीति प्रस्तुत करते हैं जो वितरित कूलिंग, घटते ओवरहेड्स को सक्षम बनाता है। इसके अलावा, हम कॉमिट, एकीकृत कोर और मेमोरी थर्मल सिमुलेशन टूलचेन पेश करते हैं, जो विभिन्न कोर और तीन-डी मेमोरी कॉन्फिगरेशन का समर्थन करता है। कॉमिट के स्रोत कोड को एम.आई.टी लाइसेंस के तहत सार्वजनिक उपयोग के लिए खुला कर दिया गया है।

Contents

Certificate	i
Acknowledgments	iii
Abstract	v
List of Figures	xvi
List of Tables	xvii
List of Acronyms	xix
1 Introduction	1
1.1 3D Memory Architectures	2
1.1.1 Industry Standard 3D TSV Memory Architectures	3
1.1.2 3D Memory Challenges	4
1.2 Thermal Issues in Processors	6
1.3 Background and Related work	7
1.3.1 Simulation Tools and Models	7
1.3.2 Power Management	9
1.3.3 Thermal Management	11
1.3.4 Limitations of Prior Work	15
1.4 Summary and Outline of the Thesis	16
2 Leakage Aware Thermal Management for 3D Memories	19
2.1 Background and Motivation	21
2.2 The proposed DTM strategies	22
2.2.1 Hybrid 3D + 2D Memory Architectures	22
2.2.2 TAM (Thermal-Aware Migration)	23

2.2.3	Memory Delay Model	24
2.2.4	FastCool	27
2.2.5	FC Policy Improvements	29
2.2.6	Energy-Efficient FastCool (EEFC)	31
2.2.7	Leakage Current Estimation	35
2.2.8	DTM Policy Implementation	35
2.3	Experiments and Results	37
2.3.1	Experimental Setup	37
2.3.2	Results and Discussion	41
2.3.3	Determining DTM Parameters	46
2.3.4	Sensitivity Analysis	50
2.4	Effect of Process Variations	53
2.4.1	Variability Modeling	53
2.4.2	Experimental Setup	54
2.4.3	Results and Discussion	55
2.4.4	Sensitivity Analysis	60
2.5	Summary	61
3	PredictNcool: Thermal Management Using Temperature Prediction	63
3.1	Motivation	64
3.2	Our Proposal	66
3.2.1	Design of a Temperature Predictor	66
3.2.2	Design of PredictNcool (PNC)	74
3.2.3	Memory System Energy Modeling	76
3.3	Experimental Results and Analysis	78
3.3.1	Results: Temperature Predictor	79
3.3.2	Results: PredictNcool DTM policy	81
3.4	Summary	86
4	Utilizing low power states for thermal management	87
4.1	Low Power State Based DTM Polices	90
4.1.1	Rotating Channel Low Power State based DTM (RL-DTM)	90
4.1.2	Masked Rotating Channel Low Power State based DTM (ML-DTM)	91
4.2	Integrated DTM Strategy (PML-DTM)	96
4.2.1	Partial Channel Closure	97
4.2.2	PML-DTM Policy	99

4.3	Experiments and Results	100
4.3.1	Experimental Setup and Methodology	100
4.3.2	Results and Analysis	100
4.3.3	Scalability Analysis	106
4.3.4	Parameter Determination	108
4.3.5	DTM Evaluation for Multithreaded Workloads	111
4.3.6	DTM Evaluation for High Bandwidth Memory (HBM2E)	115
4.4	Summary	117
5	CoreMemDTM: Integrated Processor Core and 3D Memory DTM	119
5.1	Motivation	120
5.2	<i>CoreMemDTM</i> : Proposed DTM Approach	121
5.2.1	<i>CoreDTM</i> : Multi-level Stall-Balanced DTM for Core	122
5.2.2	<i>MemDTM</i> : Low Power State based DTM for Memory	126
5.2.3	<i>CoreMemDTM</i> : Integrated DTM for Core and Memory	127
5.3	Experiments and Results	128
5.3.1	Experimental Setup	128
5.3.2	Results and Discussion	129
5.4	Summary	134
6	CoMeT: An Integrated Interval Thermal Simulation Toolchain	135
6.1	<i>CoMeT</i> : Integrated Thermal Simulation for Cores and Memories	137
6.1.1	<i>CoMeT</i> Tool Flow	138
6.1.2	Support for Multiple Core-Memory Configurations	141
6.1.3	Leakage-Aware Thermal Simulation for Memories	142
6.1.4	Simulation Control Options	142
6.1.5	SchedAPI: Resource Management Policies for Application Scheduling, Mapping, and DVFS	143
6.1.6	<i>HeatView</i>	144
6.1.7	Floorplan Generator	145
6.1.8	Automated Build Verification (Smoke Testing)	145
6.2	Experimental Studies	146
6.2.1	Experimental Setup	146
6.2.2	Thermal Profile for Various Architecture Configurations	147
6.2.3	Thermal Profile for Various Benchmarks	150
6.2.4	Case Study: Thermal-Aware Scheduler and DVFS	151

6.2.5	Parameter Variation	153
6.2.6	Overhead Analysis	154
6.3	Summary	155
7	Conclusion and Future Directions	157
7.1	Future Work	159
	References	163
	List of Publications	183
	Biography	185

List of Figures

1.1	Modern Architectures	2
1.2	Industry Standard 3D TSV Memory Architectures	4
1.3	8-layer, 16-channel 3D memory	6
1.4	Rank (8MB) power vs. temperature for a 3D memory using CACTI-3DD [27]	6
1.5	Core thermal map (under uniform power dissipation).	7
1.6	Core power versus temperature	7
1.7	PID controller for processor thermal management	12
1.8	Thermal-aware (3D + 2D) memory management	14
1.9	Thermal-aware page allocation	15
2.1	16-core Processor with off-chip 3D and 2D memories.	20
2.2	Effect of leakage power on the 3D memory steady-state temperature	21
2.3	Thermal-aware data migration and channel shut down sequence	24
2.4	Memory temperature versus access rate	28
2.5	Insights for improving the DTM policy	30
2.6	<i>EEFC DTM policy</i>	32
2.7	Rank power versus temperature	40
2.8	Steady-state temperature for various workloads	40
2.9	Transient temperature for Icepak and HotSpot	40
2.10	Normalized execution time comparison	41
2.11	2D memory delay comparison	41
2.12	Temperature-time trace for <i>bgs1</i> workload	43
2.13	DTM state transitions for <i>bgs1</i> workload	43
2.14	Normalized memory energy (3D+2D) comparison	44
2.15	Normalized EDP comparison	44
2.16	Active time for various policies	45
2.17	Normalized dynamic memory energy comparison	46
2.18	Normalized memory EDP comparison for training workloads	47

2.19	Normalized memory EDP comparison for test workloads	47
2.20	Layer-wise temperature	49
2.21	Actual versus estimated leakage	49
2.22	Execution time comparison for 16/24/32 core architectures	51
2.23	Execution time comparison with 2/3/4 memory channels	52
2.24	Effect of variability on steady-state temperature of workloads	53
2.25	Box plot elaboration	56
2.26	Execution time comparison for $\sigma_{WID} = 9.2\%$, $\sigma_{D2D} = 5\%$	56
2.27	Temperature-time trace for <i>bgs12</i>	57
2.28	Energy for different workloads using EEFC	58
2.29	EDP for different workloads using EEFC	59
2.30	Benefits of considering access rate for EEFC	59
2.31	Normalized EDP for LBM	61
2.32	Benefits of considering temperature for EEFC	61
3.1	Linear steady-state temperature prediction for a 3D memory.	65
3.2	Reactive versus proactive DTM strategy for bwaves.	65
3.3	Steady-state temperature predictor for 3D memory	68
3.4	Bottom layer weights	68
3.5	Symmetric ranks with respect to a rank R	70
3.6	Bottom layer weights using 8M datapoints	70
3.7	Symmetrically located ranks in a layer.	71
3.8	Bottom layer weights for rank (2,1,1)	71
3.9	Base layer steady-state temperature prediction	73
3.10	Final weights learned for a rank R	78
3.11	Linear temperature predictor based on symmetry	79
3.12	<i>sopl</i> temperature trace	82
3.13	Runtime comparison	84
3.14	Energy comparison	84
3.15	EDP comparison	84
3.16	Runtime comparison (for varying <i>THR</i>)	85
3.17	Energy comparison (for varying <i>THR</i>)	85
4.1	Rotating power states for 8 channels	88
4.2	Partial channel closure	88
4.3	Fine grained rotating channel DTM policy	92

4.4	A rotating channel DTM policy	96
4.5	Thermal profile (in°C)	97
4.6	PML-DTM: Integrated DTM strategy	100
4.7	Execution time normalized to baseline policy	101
4.8	Execution time normalized to the No Thermal constraint condition	101
4.9	<i>lbm</i> workload temperature-time trace	103
4.10	Energy consumption normalized to baseline policy	104
4.11	Energy consumption normalized to the No Thermal constraint condition	104
4.12	EDP normalized to baseline policy	105
4.13	EDP normalized to the No Thermal constraint condition	105
4.14	Execution time (32 Core)	107
4.15	Energy consumption (32 Core)	107
4.16	EDP (32 Core)	107
4.17	Normalized execution time (w.r.t. PML-DTM)	110
4.18	Normalized EDP (w.r.t. ML-DTM)	110
4.19	Execution time for multithreaded workloads (normalized to FC)	113
4.20	Memory energy consumption for multithreaded workloads (normalized to FC)	113
4.21	EDP for multithreaded workloads (normalized to FC)	113
4.22	Execution time under uniform accesses (normalized to FC)	114
4.23	Execution time for multi-threaded workloads with HBM2E (normalized to FC)	116
4.24	Memory energy for multi-threaded workloads with HBM2E (normalized to FC)	116
4.25	EDP for multi-threaded workloads with HBM2E (normalized to FC)	116
5.1	Processor and 3D memory with thermal limits	120
5.2	Thermal map for baseline DTM policy	121
5.3	<i>CoreDTM</i> Policy with coarse grained DVFS and fine grained stall balancing	122
5.4	<i>CoreDTM</i> illustration	123
5.5	<i>MemDTM</i> illustration	126
5.6	Execution time normalized to baseline policy	130
5.7	Energy consumption normalized to baseline policy	131
5.8	Temperature-time trace for <i>mcf</i>	132
5.9	Execution time normalized to no thermal condition	132
5.10	Energy consumption normalized to no thermal condition	133
5.11	Breakup of energy consumption for memory and core	134
6.1	Various core-memory configurations	136

6.2	Overview of <i>CoMeT</i> Flow	139
6.3	<i>CoMeT</i> Detailed Flow	140
6.4	<i>SimulationControl</i>	142
6.5	Thermal profile for <i>2D-ext</i> core-memory configuration	148
6.6	Thermal profile for <i>3D-ext</i> core-memory configuration	148
6.7	Thermal profile for <i>2.5D</i> core-memory configuration	149
6.8	Thermal profile for <i>3D-stacked</i> core-memory configuration	149
6.9	Detailed/2D view of each layer	150
6.10	Temperature for three different benchmark suites	151
6.11	3D architecture transient temperature	152
6.12	Normalized speed-up and increase in temperature	153
6.13	Thermal profile of 2 layer core and 8 layer memory	154
6.14	Simulation time normalized to HotSniper toolchain.	155

List of Tables

1.1	Prior arts in simulation tools and models	8
1.2	Prior art in power and thermal management of processor and memories	10
1.3	Proposed policies and tools	17
2.1	Parameter definitions	26
2.2	Different DTM states in EEFC	31
2.3	Core architecture parameters	37
2.4	HotSpot configuration values	38
2.5	Workloads and benchmark details	38
2.6	Workloads and benchmarks for studying the effect of process variations	54
3.1	Properties of learned weights with varying database size	69
3.2	Summary of the steady-state temperature predictors for 3D memories	73
3.3	Parameter definitions for energy modeling	76
3.4	Comparison of proposed and simple steady-state temperature predictors	78
3.5	Proposed predictor characteristics for various 3D memory capacities	81
4.1	Proposed DTM polices	89
4.2	DTM states in partial channel closure	98
4.3	Access rate versus number of bottom ranks closed	98
4.4	Various multithreaded workloads	111
5.1	<i>CoreMemDTM</i> Policy	127
5.2	Workloads and benchmarks	129
6.1	Core and memory parameters	147
6.2	List of benchmarks	147

List of Acronyms

API	Application programming interface
CMPs	Chip multi-processors
DRAM	Dynamic random access memory
DTM	Dynamic thermal management
DVFS	Dynamic voltage and frequency scaling
EDA	Electronic design automation
EDP	Energy-delay product
EEFC	Energy-Efficient FastCool
FC	FastCool
HBM	High Bandwidth Memory
HMC	Hybrid Memory Cube
JEDEC	Joint Electron Device Engineering Council
LR	Linear regression
MPSoC	Multi-Processor System on Chip
NN	Neural network
NVM	Non-Volatile memory
PCB	Printed circuit board
PCM	Phase change memory
PID	Proportional-Integral-Derivative
PNC	PredictNcool
TAM	Thermal-aware migration
WID	Within-die
WIO	Wide I/O

