

APPLICATION-AWARE DYNAMIC THERMAL MANAGEMENT IN 3D DRAM

SHAILJA PANDEY



**DEPARTMENT OF COMPUTER SCIENCE & ENGINEERING
INDIAN INSTITUTE OF TECHNOLOGY DELHI
JULY 2024**

©Indian Institute of Technology Delhi (IITD), New Delhi, 2024

APPLICATION-AWARE DYNAMIC THERMAL MANAGEMENT IN 3D DRAM

by

SHAILJA PANDEY

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

JULY 2024

To
Amma, Kakka, Mummy, and Papa,
who blessed me with Education.

Certificate

This is to certify that the thesis titled “**Application-aware Dynamic Thermal Management in 3D DRAM**” being submitted by **Shailja Pandey** for the award of **Doctor of Philosophy** in Computer Science and Engineering is a record of bonafide work carried out by her 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 with other students, which have been explicitly specified/acknowledged in the corresponding chapters and the part done by those collaborators appeared in their respective reports/theses.

Preeti Ranjan Panda

Professor

Department of Computer Science and Engg.

Indian Institute of Technology Delhi

New Delhi- 110016

Acknowledgments

“You are what you believe in. You become that which you believe you can become.”

—Bhagavad Gita

I would begin by expressing my deepest gratitude to my Ph.D. supervisor, Prof. Preeti Ranjan Panda, for his patient guidance, enthusiastic encouragement, and invaluable pieces of advice throughout my Ph.D. journey. I joined Prof. Panda’s group in February 2020 after three and a half years into the Ph.D. program. Shortly after that, there was a nationwide COVID lockdown while I was restarting my research in a different area with no prior research experience in the area. That was the toughest time in this whole journey as I was navigating through three major tasks: (1) motivate myself to continue my Ph.D., (2) learn a new research area from scratch and define an interesting problem statement, and (3) deal with the anxiety and uncertainty that COVID brought into the world. Amid all the chaos, Prof. Panda showed me a path that I eventually started believing would converge in the future. His expertise and insights were instrumental in shaping my research ideas and methodologies. He is the best mentor anyone can ask for. During my interactions with him, I have learned so many subtle ways of mentoring a student that I will try to replicate when the opportunity comes. As we say among ourselves in the lab – we go into his office with so much stress and come out with so much motivation. Most importantly, I will remain indebted to him for teaching me so much beyond technical lessons and research, indirectly influencing my personal growth as well.

My heartfelt thanks to the members of my Ph.D. committee, Prof. Kolin Paul, Prof. Anuj Grover (IIIT Delhi), and Prof. Abhilash Jindal, for their constructive feedback and suggestions in the semester’s progress presentations. Their diverse perspectives improved my research proposals and contributed significantly to their quality.

When Prof. Panda and I defined our problem statement around deep neural networks (DNNs), my friend, Vishal Sharma, helped me understand the fundamentals of DNNs and pointed me to key resources for making faster progress. I am extremely thankful and indebted to him for those countless discussions on how to design efficient Systems for ML workloads. Several insights used in this thesis came from those discussions and helped me build the foundations of my studies.

I would like to express my sincere thanks to everyone I have had the chance to collaborate with in different ways and learn so much along the process. I worked with Lokesh Siddhu, Prof. Rajesh Kedia (IIT Hyderabad), Prof. Anuj Pathania (University of Amsterdam), Martin Rapp, and Prof. Jörg Henkel (Karlsruhe Institute of Technology) for the development of the

simulation infrastructure used in this research. Along with Lokesh, I briefly collaborated with Aritra Bagchi, Issar Ahmad, and Rajesh in one of the research studies. Working with BTech students – Mritunjay Kumar Gupta (helped in modeling HBM architecture), Sayam Sethi (worked on 3D DRAM power budgeting problem), and Aaveg Jain (helped with 12-layer 3D DRAM study) was very interesting as it helped me develop mentoring skills at some level. Overall, these collaborations and interactions contributed greatly to my overall research experience, making it more enriching and rewarding.

My heartfelt thanks to all my friends in our Memory and Embedded Architecture (MARG) Group – Rajesh Kedia, Lokesh Siddhu, Aritra Bagchi, Dharamjeet Choudhary, Sakshi Tiwari, Garima Modi, Divya Praneetha, Ayushi Agarwal, Naman Jain, Dinesh Joshi, Radhika Dharwadkar, and Ravikant Bhardwaj. Lokesh helped and guided me with the initial infrastructure setup required for the research and contributed to the brainstorming sessions to define the problem statement. While Vishal immensely supported me throughout the process when I was looking to change my Ph.D. supervisor, Rajesh and Lokesh gave key suggestions that helped me ground my thoughts for approaching Prof. Panda as a potential advisor. Thank you to all of you for the lovely memories in the lab, *gupshup* over tea and snacks, long chit-chats in the corridors, and fun-filled group lunches. The friendly and positive lab environment greatly helped me settle down quickly and accelerate my research work. Thanks to Sanjana for the joyful time in the initial years, and to Nikki for giving a nice company in the lab. I will deeply cherish all the memories with each one of you for the rest of my life.

I am sincerely thankful to the office staff of the CSE department, particularly Vandana ma'am, Rekha ma'am, Hemant sir, Rajesh sir, and Suresh sir, for their timely efforts toward several administrative tasks involved in my Ph.D. journey. I remember the joyful conversations with Vandana Ma'am that provided me with lots of encouragement when I shifted to a new lab during COVID and was alone in the lab for several months. Thank you to the cleaning staff and guards in the SIT building who spread positivity in their own way by smiling on seeing the students and participating in casual conversations.

I am grateful to have spent several years on the IIT campus and made so many loving memories that will forever remain with me and make me smile. Some of my fondest memories at IIT Delhi during the seven years of stay on campus are in the Kailash Hostel. Thanks to my roommate Khushboo for being there during good and bad times. I would also like to thank two of my old friends, Neetika Singh and Arpita Shah, who supported me during various high and low points in my life.

I would like to express my gratitude to my institute, IIT Delhi, for giving me the opportunity to pursue a Ph.D. and experience the whole process of discovering myself in a way that would

otherwise not have been possible. It gives me immense pride to be a part of an institute that has such deep roots. I would also like to thank the Ministry of Education of India for providing me with a scholarship during my Ph.D. and the Semiconductor Research Corporation for providing me with funding in the later years.

While one can never thank one's parents enough for nurturing them and shaping their lives in the best possible way, I am eternally grateful to my parents, Maya and Banspati Pandey, for providing me with their unwavering support during this journey and breaking several societal norms for me that they grew up with and gotten old with. My brothers, Ankur and Anupam, have been my biggest cheerleaders throughout my education journey, at times prioritizing my seamless education over theirs and trying to eliminate all the hurdles. During my Ph.D., my brothers got married to my lovely sisters-in-law – Kamana Bhabhi and Deepti Bhabhi, who brought so much happiness to the family. Further, both my brothers were blessed with their daughters – Aarya and Aadvika, making me a proud *Bua*. My heart is filled with so much love for my nieces. I would like to thank my *naani* for being a great example of a courageous woman for me. I am also grateful to my *foofaji*, Late Krishna Kishore Shukla, for accompanying me during interviews at various IITs and otherwise while I had still not learned to travel alone.

Finally, I would like to acknowledge that it is a huge privilege to be born in India and, more so, in a village of India that provides an unparalleled life experience. I am thankful to our goddess *Maa Sharda* for blessing me to be born at the epicenter of freedom and spirituality. *Maa Sharda* is my primary source of strength and inspiration in easy and tough times, and I owe everything to her.

Shailja Pandey

Abstract

3D DRAM is a promising technology to meet the bandwidth demands of data-hungry ML/AI applications and is being actively used in a wide range of hardware platforms such as GPUs, FPGAs, DNN accelerators, and general-purpose cores. Despite the potential of 3D DRAM to bridge the memory bandwidth gaps, the vertically stacked integration of DRAM dies brings a new set of system design challenges. One such challenge is the poor heat dissipation capability, causing excessive memory heating and eventually throttling the memory, slowing down the applications.

Heating in 3D DRAM, due to high power density and close proximity of memory cells in stacked architectures, impacts the data retention capabilities, increasing the refresh rates significantly. 3D stacking enables hundreds of banks/ranks and several independent channels per memory device such that each channel is connected to a separate memory controller; however, the off-chip memory bandwidth is often limited due to limited power budgets that might arise in practical scenarios. Due to this, it is often not possible to activate and run all memory ranks simultaneously, similar to *dark silicon* in multi-core systems. Systems running under a fixed power budget and a thermal constraint employ dynamic power budgeting (*DPB*) and dynamic thermal management (*DTM*) policies, resulting in performance and energy overheads due to thermal-induced CPU and memory throttling.

Existing solutions towards dynamic thermal management are application-agnostic and often lead to sub-optimal performances. Modern applications execute in distinct *phases*, exhibiting varying memory access behavior and showing different sensitivities to architectural parameters such as last level cache (LLC) size, prefetch settings, memory-level parallelism, and memory bandwidth. The indiscriminate approach of prior works for DTM in 3D DRAM may lead to avoidable performance degradation, as different application phases have varying contributions to memory temperature. To this end, this dissertation proposes application-aware DTM techniques targeted towards minimizing the associated performance penalties. We present novel micro-architecture-based approaches and system-level policies for managing heat in 3D DRAM and facilitate faster cooling, leveraging knowledge of the workload. We

focus on deep neural networks (DNNs) in this thesis.

Deep neural network (DNN) implementations are typically characterized by huge data sets and concurrent computation, resulting in a demand for high memory bandwidth due to intensive data movement between compute units and off-chip memory. Performing DNN inference on a general-purpose processor is common, and multi-core CPUs are being used for DNN implementation in servers and commercial SoCs. 3D DRAM is becoming a key enabler for high memory-intensive applications such as DNNs running on general-purpose cores. This dissertation focuses on efficient thermal management on a system consisting of general-purpose CPU cores and a thermally-constrained 3D DRAM.

Firstly, we present a memory temperature-aware prefetcher that jointly utilizes DNN application’s memory behavior and 3D DRAM temperature to minimize DTM overheads. The customized prefetching mechanism dynamically computes the best *prefetch settings* for different DNN workload phases, utilizing the unused 3D DRAM bandwidth without adversely affecting the workload’s thermal footprint.

Secondly, we present an efficient task mapping based DTM policy for multi-core processors to minimize memory throttling through proactive management of the potential thermal hotspots in 3D DRAM using application-aware optimizations in processing cores. An application-aware DVFS is used to reduce the thermal impact of the application’s *phase* when task mapping to an ideal core is not feasible.

Thirdly, we present a thermal and memory access pattern-aware data mapping based DTM policy to map a DNN layer’s data to 3D DRAM pseudo-channels. An intelligent mapping of the DNN application’s data across memory channels, pseudo-channels, and banks, leveraging the knowledge about each layer’s memory access pattern and data reuse, significantly reduces the effective memory latency and the thermal-induced application slowdown.

Finally, we present a reward-based and adjacency-aware dynamic power budgeting which periodically suggests the ideal set of channels that should be enabled under a given power budget and thermal constraint, performing a coordinated thermal and power management for 3D DRAM. The thermal characteristics of 3D DRAM make power budgeting a non-trivial task that requires run-time intervention to minimize the associated performance penalties.

To evaluate the proposed policies, an integrated performance-thermal simulation framework is used. Our results demonstrate significant DNN inference time and memory energy reductions over the state-of-the-art, with minimal run-time cost. This dissertation establishes the significance of leveraging application knowledge towards efficient thermal management in 3D DRAM for a class of well-structured workloads such as deep neural networks.

- / ö , â õ l
ö 3 ¼ , â l â
l , ö â â , ñ
ç ñ â l â ,
¼ / , ¼ , ñ
; , - â , ý - l ,
ÿ ñ õ â l
l , l - l
l
l ö - ½ â - l
ÿ l ö - ä â , - â
l â -
l , â õ
ö l - ,
¼ ö - l â -
- ö ½ ö
ö l ö ö
ö l () ö -
 , ö ç l
é ö ý -
ö ç - ö â
l ö l - â
l
 , - l
ö ÿ ö â
ä , ö l ö â

ö - ,
ô - ý
| |
ö - ,
- |
â - †
- |
ö ½ , ñ | -
- † ö -
¼ | |
- ç |
| - ¼
ý - |
ñ - ñ |
ö | ç
ñ | ½

Contents

Certificate	i
Acknowledgments	iii
Abstract	vii
List of Figures	xx
List of Tables	xxi
1 Introduction	1
1.1 Introduction to 3D DRAM	3
1.1.1 Thermal Characteristics of 3D DRAM	4
1.2 Thermal Mitigation in 3D DRAMs	5
1.2.1 Dynamic Thermal Management	6
1.2.2 Proactive vs. Reactive Approaches	6
1.2.3 Application-aware Dynamic Thermal Management	6
1.3 Deep Neural Networks	7
1.3.1 Hardware for DNN Inference	9
1.3.2 Thermal Footprint of DNN on 3D DRAM	11
1.4 Contributions of the Thesis	12
1.5 Dissertation Outline	13
2 System Architecture and Modeling Framework	15
2.1 System Architecture	15
2.2 Experimental Framework	15
2.3 System Models	16
2.3.1 3D DRAM Power Model	16
2.3.2 3D DRAM Thermal Model	17

2.3.3	Core-to-channel Mapping	18
2.3.4	System Constraints	19
3	Thermal-aware Customized Prefetching	21
3.1	Novel Contributions	23
3.2	The NeuroCool Approach	23
3.2.1	Low-power States Based DTM	23
3.2.2	Customized Prefetching	24
3.2.3	Thermal Impact of Data Reuse Policies	26
3.2.4	Overall Flow	28
3.2.5	Classifying DNN Layers	28
3.2.6	Phase and Memory Temperature-Aware Customized Prefetching (TAP)	30
3.2.7	Dynamic Thermal Management	32
3.2.8	The NeuroCool DTM Policy	36
3.3	Experimental Evaluation	37
3.3.1	Simulation Environment	37
3.3.2	Benchmarks Used	38
3.3.3	AlexNet Characterization	39
3.3.4	Performance & Energy Improvement	41
3.3.5	Transient Temperature Behavior	42
3.3.6	Parameter Selection	45
3.3.7	Sensitivity Analysis	46
3.3.8	Transformer-based Models	50
3.3.9	Overhead Analysis	54
3.4	Summary	55
4	Thermal-aware Task Mapping	57
4.1	Motivation	57
4.2	Novel Contributions	61
4.3	The NeuroMap Approach	62
4.3.1	Problem Definition	62
4.3.2	Proposal Overview	63
4.3.3	Static Core to HBM Channel Mapping	64
4.3.4	Application-Aware Task-to-Core Mapping	65
4.3.5	Application Swapping and Application-aware DVFS	66
4.3.6	Low-power States Based DTM on Memory Side	67

4.3.7	The NeuroMap Policy	69
4.3.8	Implementation And Efficiency	70
4.4	Experimental Evaluation	72
4.4.1	Simulation Environment and Workloads	73
4.4.2	Baselines	74
4.4.3	Performance and Memory Energy Improvement	74
4.4.4	Performance and Energy Break-up of Optimizations	75
4.4.5	Transient Temperature Behavior	76
4.4.6	Order of Invocation of Optimizations	77
4.4.7	Sensitivity to L3 Cache Size	78
4.4.8	Overhead Analysis	79
4.5	Summary	80
5	Thermal-aware Data Mapping	83
5.1	Motivation	86
5.2	Novel Contributions	89
5.3	The NeuroTAP Proposal	90
5.3.1	Offline Characterization of Deep Neural Networks	91
5.3.2	Thermal and Application-aware Data Mapping in 3D DRAM	94
5.3.3	DRAM Low-power States Based Dynamic Thermal Management	98
5.3.4	The Overall Flow	99
5.4	Experimental Evaluation	101
5.4.1	Simulation Workloads	102
5.4.2	Baselines	103
5.4.3	Inference Time and Memory Energy Improvement	104
5.4.4	Thermal Profile of 3D DRAM using DTM Policies	106
5.4.5	Impact of Number of Cores & Last-level Cache Sizes	107
5.4.6	Applicability to Transformer-based Models	108
5.4.7	Implementation Aspects	110
5.5	Summary	112
6	A Heuristic-based Dynamic Power Budgeting	115
6.1	Motivation	117
6.2	Novel Contributions	120
6.3	3D DRAM Power Budgeting and Thermal Management	120
6.3.1	Problem Definition	120

6.3.2	Preliminary Power Budgeting Policies	121
6.4	The 3D-TemPo Approach	124
6.4.1	Reward Computation	125
6.4.2	Incorporating Adjacency-awareness	125
6.4.3	DRAM Low-power Based DTM	127
6.4.4	The Overall Flow	127
6.5	Experimental Evaluation	128
6.5.1	Simulation Environment	128
6.5.2	Performance Improvement	130
6.5.3	Memory Energy Improvement	132
6.5.4	Analysis of Policy Behavior	134
6.5.5	Transient Temperature Behavior	135
6.5.6	Discussion on Working of Policies	136
6.5.7	Implementation Details	138
6.6	Summary	138
7	Conclusion	141
7.1	Future Research Directions	143
	Bibliography	145
	A Notations	161
	B List of Acronyms	163
	C Detailed Acknowledgements for Collaborative Works	165
	List of Publications	167
	Biography	169

List of Figures

1.1	A second generation HBM2E memory architecture. The <i>E</i> stands for <i>Evolutionary</i>	3
1.2	A typical deep neural network (DNN) architecture.	8
1.3	Variation in the number of DRAM accesses per milli-second across different layers of AlexNet. "E" represents the end of previous layer and the start of next layer.	9
1.4	Variation DRAM access counts across different Alexnet layers with different data reuse policies — Input, Weight, Output.	10
1.5	Transient temperature of 3D DRAM running DNNs in absence of dynamic thermal management (DTM). <i>CE</i> and <i>PE</i> mark the end of convolution layer and pooling layer respectively.	11
2.1	System architecture comprising general-purpose cores and off-chip 3D DRAM.	16
2.2	Interface between processing cores and HBM2E.	17
2.3	Placement of temperature sensors in DRAM dies of HBM2E. The label x.y.z in the figure indicates (stack_id).(dram_die_number).(channel_number).	18
3.1	Prefetch settings: <i>degree</i> , <i>interval</i> and <i>stride</i> . <i>Block</i> corresponds to a 64 Bytes chunk inside a 4KB physical frame/page in memory.	22
3.2	Variation in MPKI of DNN workloads for different L3 cache sizes.	23
3.3	Transient temperature behavior of a DNN workload with prefetching ON/OFF without a DTM.	26
3.4	Transient memory access trace and transient temperature of a DNN workload with different data reuse policies in absence of DTM.	27
3.5	NeuroCool workflow.	28
3.6	Determining optimal number of clusters.	29
3.7	Cluster assignment to DNN layers.	30

3.8	Memory state transitioning in a <i>TierWise</i> Low-power states based DTM policy. Shades of red, yellow, and blue colors indicate higher temperatures, moderate temperatures, and low temperatures respectively.	33
3.9	Working of <i>HyDTM</i> to select the granularity (ranks/tiers/DRAM) for transition to low-power states based on <i>access level parallelism</i> in the current DNN layer. Execution time normalized to <i>NoDTM</i>	34
3.10	Alexnet layer-wise performance.	40
3.11	DNN performance of training and test data with different schemes, normalized to <i>FastCool</i>	41
3.12	Memory energy of training and test data with different schemes, normalized to <i>FastCool</i>	42
3.13	Transient temperatures when running 32 instances of a VGG layer with different policies under thermal constraints. "A", "B", "C", "D", and "E" indicate completion times with <i>NeuroCool</i> , <i>DTM+DefPF</i> , <i>OnlyDTM</i> , <i>FastCool</i> , and <i>NoDTM</i> policies.	43
3.14	Normalized execution time comparison for (a) training and (b) test workloads with different <i>customized prefetching</i> temperature threshold T_{th} values.	45
3.15	DNN performance of training and test data with different schemes normalized to <i>FastCool</i> with critical temperature $T_{crit}=90^{\circ}\text{C}$	47
3.16	DNN performance of training and test data with different schemes normalized to <i>FastCool</i> with critical temperature $T_{crit}=100^{\circ}\text{C}$	48
3.17	Comparison of normalized execution time on a 64-core system at different critical temperatures.	49
3.18	Normalized execution time with 3D DRAM consisting of 12 DRAM tiers. $T_{crit} = 80^{\circ}\text{C}$	51
3.19	Architecture of BERT Model. (i) Overall Structure, (ii) Transformer Encoder Block, (iii) Multi-head Attention Block, and (iv) Attention Layer.	52
3.20	Normalized execution time of BERT _{BASE}	53
3.21	Architecture of a Decoder-only Transformer Model.	54
3.22	Normalized execution time of GPT-2.	54
4.1	Mapping decisions of 4 task mapping policies. {1,2,3,4} represents 4 different applications running on a 16 core system.	58
4.2	Variation in DRAM access counts of AlexNet with time using different data reuse policies — Input, Weight, Output.	59

4.3	Transient temperatures of a DNN workload (AlexNet + ResNet18) with four mapping policies under thermal constraints. "A", "B", "C", and "D" indicate completion times with <i>ApplicationAware</i> , <i>RoundRobin</i> , <i>Random</i> , and <i>FirstUnused</i> policies.	60
4.4	Transient average temperature of a DNN workload (AlexNet + ResNet18) with four mapping policies under thermal constraints. "A", "B", "C", and "D" indicate completion times with <i>ApplicationAware</i> , <i>RoundRobin</i> , <i>Random</i> , and <i>FirstUnused</i> policies.	61
4.5	Transient maximum temperature gradient of a DNN workload (AlexNet + ResNet18) with four mapping policies under thermal constraints. "A", "B", "C", and "D" indicate completion times with <i>ApplicationAware</i> , <i>RoundRobin</i> , <i>Random</i> , and <i>FirstUnused</i>	62
4.6	The <i>NeuroMap</i> scheme.	63
4.7	Static core to channel mapping to minimize vertical stacking of thermal hotspots in 3D DRAM.	64
4.8	DNN layer's memory access characteristics aware task-to-core mapping. . . .	66
4.9	Swapping applications across core groups when ideal core is unavailable. . . .	66
4.10	Application-aware DVFS on current core when potential swap candidates are unavailable.	67
4.11	Memory channel state transitioning in low-power states based DTM policy. . .	68
4.12	Flowchart of task-to-core mapping.	72
4.13	Execution time of DNN workloads with six policies.	75
4.14	Memory energy of different DNN workloads with different policies normalized to <i>FastCool</i>	75
4.15	Impact of enabling/disabling individual optimizations on performance and energy (normalized to <i>LPS</i> optimization).	76
4.16	Transient temperatures of WK-2.	77
4.17	Transient temperatures of WK-4.	78
4.18	<i>Application-aware DVFS</i> vs. <i>Application Swapping</i>	79
4.19	Performance sensitivity to L3 cache sizes.	80
4.20	Transient memory access trace of WK-2 with <i>NeuroMap</i> policy depicting task migration, application swapping, and application-aware DVFS during its execution.	81

5.1	A typical DNN architecture and the reuse policies for three datatypes – <i>weight</i> , <i>input</i> , and <i>output</i> . In (a), (b), and (c), the value/pixel shown in <i>green</i> color is <i>reused</i> in computation across convolutional operations for all the values/pixels shaded in <i>blue</i> color.	84
5.2	Reuse of three datatypes – input, output, and weight in different AlexNet layers.	85
5.3	Reuse-per-kilo-instructions in different AlexNet layers.	86
5.4	Peak 3D DRAM temperature while running AlexNet <i>CONV2</i> layer, exhibiting <i>strided</i> memory access pattern, mapped to pseudo-channels located in <i>top</i> , <i>middle</i> , and <i>bottom</i> dies.	87
5.5	Peak 3D DRAM temperature while running AlexNet <i>FC7</i> layer, exhibiting <i>sequential</i> memory access patterns, mapped to pseudo-channels located in <i>top</i> , <i>middle</i> , and <i>bottom</i> dies.	88
5.6	Impact of data mapping to <i>top</i> , <i>middle</i> , and <i>bottom</i> pseudo-channels on inference time of AlexNet <i>FC7</i>	89
5.7	Overall flow of the proposed approach.	90
5.8	Thermal-aware mapping of DNN layers to 3D DRAM pseudo-channels . (<i>PC:x-y-z</i>): <i>x</i> is the stack ID, <i>y</i> is the channel number within the stack, and <i>z</i> is the pseudo-channel number within the channel.	95
5.9	Working of the DRAM low-power states based DTM at a pseudo-channel level.	99
5.10	An example of <i>NeuroTAP</i> data mapping in 3D DRAM.	101
5.11	A comparison of inference time of DNN workloads with different policies, normalized to <i>NoDTM</i>	104
5.12	A comparison of memory energy of DNN workloads with different policies, normalized to <i>NoDTM</i>	105
5.13	(a) Temperature variation, and (b) number of <i>active</i> pseudo-channels over time in WK-7 for different DTM policies under a thermal limit $T_{crit}=80^{\circ}\text{C}$. "A", "B", "C", and "D" mark the completion times with <i>NoDTM</i> , <i>NeuroTAP</i> , <i>DASC</i> , and <i>NeuroMap</i> respectively.	106
5.14	Sensitivity of the <i>NeuroTAP</i> policy to L3 cache sizes.	108
5.15	Comparison of inference time with BERT _{BASE} in the mix, normalized to <i>NoDTM</i>	111
5.16	Comparison of memory energy with BERT _{BASE} in the mix, normalized to <i>NoDTM</i> .	111

5.17	Peak temperature and execution time comparison for WK-7 with enabling/disabling DVFS, and stalling the DNNs in the <i>NeuroTAP</i> policy. "A", "B", and "C" mark the completion times with <i>EnableDVFS</i> , <i>NoDVFS</i> , and <i>StallApplication</i> . □ denote the unavailability of an ideal core for one or more DNN layers.	111
6.1	Dynamic power budgeting (DPB) and dynamic thermal management (DTM) working in isolation causes contradictory decisions, over-compensation, and complete memory throttling. The colors represent the channel temperatures – (1) shades of <i>red</i> : hotter channels, (2) shades of <i>yellow</i> : moderate temperatures, and (3) shades of <i>blue</i> : cooler channels.	119
6.2	Working of <i>RoundRobin</i> policy.	122
6.3	Working of <i>Alternation</i> policy.	123
6.4	Working of <i>MostFrequentlyUsed (MFU)</i> policy.	124
6.5	Working of <i>Adjacency-aware 3D-TemPo</i> policy in the <i>critical</i> region. <i>Shading</i> represents <i>activated</i> channels.	126
6.6	Execution time of SPEC 2017 workloads with different power budgeting policies for $P_b=64\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to no power and thermal constraints (NoCons).	130
6.7	Execution time of PARSEC 2.1 workloads with different power budgeting policies for $P_b=64\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to no power and thermal constraints (NoCons).	131
6.8	Execution time of SPEC 2017 workloads with different power budgeting policies for $P_b=96\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to no power and thermal constraints (NoCons).	131
6.9	Execution time of PARSEC 2.1 workloads with different power budgeting policies for $P_b=96\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to no power and thermal constraints (NoCons).	131
6.10	Memory energy consumption of SPEC 2017 workloads with different power budgeting policies for $P_b=64\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to <i>NoCons</i>	132
6.11	Memory energy consumption of PARSEC 2.1 workloads with different power budgeting policies for $P_b=64\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to <i>NoCons</i>	133
6.12	Memory energy consumption of SPEC 2017 workloads with different power budgeting policies for $P_b=96\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to <i>NoCons</i>	133
6.13	Memory energy consumption of PARSEC 2.1 workloads with different power budgeting policies for $P_b=96\text{W}$ and $T_{crit}=80^\circ\text{C}$, normalized to <i>NoCons</i>	133

6.14	Transient temperature for <i>WK-4</i> with different power budgeting policies for $P_b=64\text{W}$	134
6.15	Comparison of (a) number of channels in ON state, (b) power budget utilization, and (c) peak memory temperature over time in different power budgeting policies.	137

List of Tables

3.1	Single-core DNN execution (prefetching ON/OFF).	25
3.2	Prefetch Class and Associated Settings for the Temperature-Aware Prefetching (TAP) policy.	31
3.3	Simulation parameters.	37
3.4	Configuration parameters of HotSpot.	38
3.5	Workload characteristics.	39
3.6	<i>NeuroCool</i> settings for Alexnet.	39
3.7	Number of <i>Thermal Stalls</i> and <i>Average Cooldown Time</i>	44
3.8	DNN layer <i>Access Level Parallelism</i> threshold.	46
3.9	Simulation workloads on a 64-core system.	48
3.10	Prefetch settings for Bert _{BASE}	53
4.1	DNN characteristics.	73
4.2	Simulation workloads.	73
5.1	DNN layer-aware application-to-pseudo-channel mapping in 3D DRAM. . . .	97
5.2	Deep neural networks used. <i>MPKI</i> denotes last-level cache misses-per-kilo-instructions, <i>RBW</i> denotes read bandwidth consumed, <i>RBH</i> denotes row buffer hit rate, and <i>BLP</i> denotes bank-level parallelism.	102
5.3	Simulation Workloads. The column <i>Data Reuse Policies</i> indicates the different policies that are used during the workload execution in its different <i>phases</i> (DNN layers of different DNNs).	103
5.4	Simulation workloads for transformer-based model.	110
6.1	Simulation Workloads.	128
6.2	Comparative analysis for different power budgets using <i>3D-TemPo</i> policy. Execution time and memory energy are normalized to $P_b = 96W$	135
6.3	Number of Thermal Stalls and Average Cooldown Time for WK-5.	136

