

SYSTEM SOFTWARE IN EDGE GPU PLATFORMS FOR URBAN SUSTAINABILITY APPLICATIONS IN DEVELOPING COUNTRIES

OMAIS SHAFI



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

March 2025

SYSTEM SOFTWARE IN EDGE GPU PLATFORMS FOR URBAN SUSTAINABILITY APPLICATIONS IN DEVELOPING COUNTRIES

by

OMAIS SHAFI

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

March 2025

To Abu, Mummy, Ubaid baya
Thanks to their unwavering support, patience, and belief in me.

Certificate

This is to certify that the thesis titled **System Software in Edge GPU Platforms for Urban Sustainability Applications in Developing Countries**, submitted by **Mr. Omais Shafi**, to the Indian Institute of Technology, Delhi, for the award of the degree of **Doctor of Philosophy** in Computer Science & Engineering, is a record of bonafide work carried out by him under my guidance and supervision at the Department of Computer Science & Engineering, Indian Institute of Technology Delhi. The contents of this thesis, in full or in parts, have not been submitted elsewhere for the award of any degree or diploma.

Dr Rijurekha Sen

Assistant Professor

Dept. of Computer Science and Engineering

Indian Institute of Technology Delhi

New Delhi-110016

Acknowledgements

For indeed, with hardship [will be] ease.

– Quran 94:5/6

All praise is due to Allah, the Almighty, upon whom we depend for sustenance and guidance. I am deeply grateful to Him for granting me the resilience and fortitude needed to persevere until the very end.

I am profoundly thankful to my advisor, Dr.Rijurekha Sen, for her timely and unwavering support, both personally and professionally. Throughout the most challenging periods of my Ph.D. journey, particularly when I faced significant obstacles and uncertainty, her steadfast assistance and encouragement were invaluable. Dr. Sen provided not only insightful guidance in my research but also the motivation and moral support needed to persevere. During times of family commitments or work obligations, her patience and constant push for my Ph.D. work were invaluable. Thank you Riju for being a *great mentor*.

I would also like to express my gratitude to my collaborators, Dr.Gayathri Ananthanarayanan, Dr.Aastha Mehta, Dr.Arpan Gujarati, Dr.Khalid Pandit, and Sachin for their valuable discussions and feedback. Regular meetings with them at various stages of my Ph.D. were crucial. I extend my thanks to my SRC members— Prof.Anshul Kumar, Prof.Kolin Paul, and Prof.Sumantra Dutta Roy—for serving on my committee and providing valuable feedback. Additionally, I am grateful to Prof.Preeti Ranjan Panda for his support during tough times and to Prof.Smruti Ranjan Sarangi for his assistance in the early years of my Ph.D. Thanks also to Chinmai and Amarjeet for their help with many experimental results.

The journey of a Ph.D. cannot be completed without a good circle of friends, and I was fortunate to have many. Mehreen deserves a special mention for her constant motivation, companionship, and all the tea and delicious food we shared. She witnessed my ups and downs and supported me like a rock. She has been my go to person and I owe her a big thanks. I want to thank Janibul Bashir for being a wonderful mentor and brother throughout the earlier years of my

Ph.D. journey. His valuable advice and support were unwavering. I also want to acknowledge Nadeem, Danish, and Zarkab. The times spent at Nilgiri hostel discussing Ph.D. life and everyday matters were some of the best moments of my journey. Nadeem, in particular, was a constant source of support when we were roommates, always taking care of my belongings and being there for me. I would also like to thank Zamir Wani for his valuable advice throughout my Ph.D. and for pushing me through the tough phases when I needed it the most. Thanks to Asifa for all her advices and for listening to my long conversations, and for the unforgettable dinners at Nilgiri hostel and the food she prepared for us.

I have no words to thank Sandeep for his support during my Ph.D. He is one of the best human beings I have come across and always pushed me forward. I will never forget our meetings at CCD, the Lebanese chicken lunches at the SIT faculty lounge, and all the outings we had together. Thanks also to Vinayak and Dishant for the wonderful outings and times shared with Sandeep. I am thankful to Ismi for being there for all the good discussions and motivation she used to give me during my Phd. Additionally, I am grateful to Shubankar, Hameedah, Diksha, and Himanshu for the enriching discussions and memorable times. I also extend my gratitude to my Kashmiri friends— Omar, Adnan, Mujeeb, Saleem, Malik, Shahid, Kaiser, Amir, Faizan, Faiz, Saqib and others. The tea and dinner parties we shared will always be cherished. I may have missed mentioning few names, but I am grateful to everyone for the cherished memories.

Imagining a journey like a Ph.D. without family support is difficult. I thank my entire family for their support and patience over these years. Thank you for your prayers and blessings. I am grateful to all my cousins for always being there when I needed them. I thank my grandparents for their blessings and support. I wish all my grandparents were here to see my growth, but I am sure those who are no longer with us are still sending their blessings.

Finally, I dedicate this thesis to three people without whom this work would never have materialized: my father, my mother, and my brother. Thank you for your unconditional support, encouragement, and the freedom you gave me at every phase of my life. I cannot possibly list everything I want to thank you for. Everything I know or have achieved is because of your efforts. My parents have gone through the toughest phases for providing me with the good education and I have no words to thank them for everything they have done for me. I owe them everything. Special mention to my brother, Ubaid Baya, who sacrificed his career by letting me study outside our hometown while he continued working there despite having job opportunities elsewhere. I owe him a lot and probably can never repay it back. Thank you, Abu, Mummy, and Ubaid Baya, for everything.

A special note of appreciation to my wife, Dr.Hadiya. Meeting you during the final stages of

my PhD has been a blessing. Your love, encouragement, and understanding have given me strength and joy.

The greatest glory in living lies not in never falling, but in rising every time we fall.

–Nelson Mandela

Omais Shafi

Abstract

Road traffic congestion increases vehicular emissions and air pollution, significantly impacting urban sustainability. The excessive pollution contributes to poor air quality and adverse health effects, while the economic costs associated with congestion and accidents strain city resources. In developing countries, where the vehicle-to-road infrastructure ratio is imbalanced, these issues are further magnified. Effective management of traffic and reduction of emissions are crucial for promoting sustainable urban development and improving the quality of life for city dwellers. Additionally, addressing these challenges is essential for reducing the carbon footprint of urban areas, mitigating climate change, and fostering resilient and livable cities. Implementing innovative traffic management solutions and investing in green infrastructure can enhance urban sustainability and create healthier, more efficient environments.

To address these challenges, advanced computational tools are being leveraged to improve road traffic measurement and management. Cameras capture video frames, which are then processed using state-of-the-art machine learning (ML) and deep neural network (DNN) models to analyze traffic patterns and enforce regulations. However, running these heavy computational models on the cloud is challenging due to the lack of reliable broadband connectivity in developing countries. An alternative is to run these models on edge devices equipped with advanced GPU platforms. This thesis explores the system software stack for edge GPU platforms, focusing on both compile-time and runtime aspects to optimize performance and manage computational demands effectively.

This thesis examines two types of system software essential for optimizing neural network inference on Nvidia edge GPUs. The first type is compile-time system software, which processes pre-trained ML models by applying various model compression techniques to generate efficient engines for inferences. The second type is runtime system software, which manages resource utilization, thermal safety, and other critical factors as the ML models perform continuous inferences. For both compile-time and runtime software, the thesis carefully audits existing software stacks from industry and academia, identifying their limitations. Based on this anal-

ysis, new software solutions are designed and implemented to address these shortcomings and enhance overall system performance. Additionally, the thesis explores various aspects of system software for Nvidia edge GPUs, including techniques for managing application dynamism and ambient temperature, and addresses the challenges of deploying applications from multiple vendors.

We first investigate key software frameworks for compressing pre-trained neural networks for edge inference, focusing on TensorRT, AutoTVM, and AutoScheduler. We develop a comprehensive set of metrics to compare these frameworks and design careful experiments to empirically measure those metrics. Our analysis of TensorRT showed notable performance and accuracy improvements but also revealed unexpected non-deterministic behaviors such as varying outputs and increased latency on more powerful hardware. We also evaluated vendor-agnostic frameworks AutoTVM and AutoScheduler, comparing them on performance predictability, inference latency, output reproducibility, and hardware utilization. This comparison led to several interesting observations which have important implications for urban sustainability applications like traffic monitoring. These insights provide recommendations for selecting the most effective DNN compiler for real-time, safety-critical applications, ensuring precise and predictable inferences.

In addition, to address the challenges of managing high throughput and ensuring system longevity in the extreme conditions of developing regions, we developed DynCNN, a controller designed for application dynamism and ambient temperature management. DynCNN leverages processor heterogeneity to regulate thread counts and accelerator frequencies, optimizing application performance while adhering to strict thermal and power constraints. We evaluated DynCNN on three commercially available embedded GPUs using a 40-day dataset from a real traffic intersection, demonstrating its effectiveness in enhancing system efficiency under these demanding conditions.

Furthermore, edge devices face challenges in deploying applications from multiple vendors, each with its own software dependencies. Using Docker containers provides isolation and manages these dependencies, enhancing security and portability. However, applications running together can affect each other's performance and increase device temperature due to shared resource contention. This necessitates advanced scheduling to ensure optimal performance and thermal management. To address these issues, we propose IceEdge, an inference serving framework for CPU-GPU based edge computing servers. IceEdge uses containers for isolation and includes a deadline- and thermal-aware scheduler for ML workloads, ensuring good GPU utilization and maintaining quality-of-service. IceEdge also incorporates a controller that precisely manages CPU-GPU operations, avoiding reliance on OS thermal control and thus

improving efficiency and stability.

In conclusion, this thesis investigates the deployment of deep neural networks (DNNs) on Nvidia edge GPU platforms for sustainability applications in developing countries. Our objective is to provide a comprehensive solution that enhances the longevity, performance, and reliability of edge computing systems by addressing critical challenges such as compiler selection, thermal management, and multi-vendor application deployment. The insights and recommendations from our research will aid in the creation of robust and efficient edge computing solutions, fostering sustainable development and improving the quality of life in underdeveloped regions.

अमूर्त

सड़क यातायात की भीड़भाड़ से वाहनों से होने वाले उत्सर्जन और वायु प्रदूषण में वृद्धि होती है, जिससे शहरी स्थिरता पर महत्वपूर्ण प्रभाव पड़ता है। अत्यधिक प्रदूषण खराब वायु गुणवत्ता और प्रतिकूल स्वास्थ्य प्रभावों में योगदान देता है, जबकि भीड़भाड़ और दुर्घटनाओं से जुड़ी आर्थिक लागत शहर के संसाधनों पर दबाव डालती है। विकासशील देशों में, जहाँ वाहन-से-सड़क अवसंरचना अनुपात असंतुलित है, ये मुद्दे और भी बढ़ जाते हैं। यातायात का प्रभावी प्रबंधन और उत्सर्जन में कमी सतत शहरी विकास को बढ़ावा देने और शहरवासियों के जीवन की गुणवत्ता में सुधार करने के लिए महत्वपूर्ण है। इसके अतिरिक्त, शहरी क्षेत्रों के कार्बन पदचिह्न को कम करने, जलवायु परिवर्तन को कम करने और लचीले और रहने योग्य शहरों को बढ़ावा देने के लिए इन चुनौतियों का समाधान करना आवश्यक है। अभिनव यातायात प्रबंधन समाधानों को लागू करना और हरित अवसंरचना में निवेश करना शहरी स्थिरता को बढ़ा सकता है और स्वस्थ, अधिक कुशल वातावरण बना सकता है।

इन चुनौतियों का समाधान करने के लिए, सड़क यातायात माप और प्रबंधन में सुधार के लिए उन्नत कम्प्यूटेशनल उपकरणों का लाभ उठाया जा रहा है। कैमरे वीडियो फ़्रेम कैप्चर करते हैं, जिन्हें ट्रैफ़िक पैटर्न का विश्लेषण करने और नियमों को लागू करने के लिए अत्याधुनिक मशीन लर्निंग (ML) और डीप न्यूरल नेटवर्क (DNN) मॉडल का उपयोग करके संसाधित किया जाता है। हालांकि, विकासशील देशों में विश्वसनीय ब्रॉडबैंड कनेक्टिविटी की कमी के कारण क्लाउड पर इन भारी कम्प्यूटेशनल मॉडल को चलाना चुनौतीपूर्ण है। एक विकल्प इन मॉडलों को उन्नत GPU प्लेटफ़ॉर्म से लैस एज डिवाइस पर चलाना है। यह थीसिस प्रदर्शन को अनुकूलित करने और कम्प्यूटेशनल मांगों को प्रभावी ढंग से प्रबंधित करने के लिए संकलन-समय और रनटाइम दोनों पहलुओं पर ध्यान केंद्रित करते हुए एज GPU प्लेटफ़ॉर्म के लिए सिस्टम सॉफ़्टवेयर स्टैक की खोज करती है।

यह थीसिस Nvidia एज GPU पर न्यूरल नेटवर्क अनुमान को अनुकूलित करने के लिए आवश्यक दो प्रकार के सिस्टम सॉफ़्टवेयर की जाँच करती है। पहला प्रकार संकलन-समय सिस्टम सॉफ़्टवेयर है, जो अनुमानों के लिए कुशल इंजन बनाने के लिए विभिन्न मॉडल संपीड़न तकनीकों को लागू करके पूर्व-प्रशिक्षित ML मॉडल को संसाधित करता है। दूसरा प्रकार रनटाइम सिस्टम सॉफ़्टवेयर है, जो संसाधन उपयोग, थर्मल सुरक्षा और अन्य महत्वपूर्ण कारकों का प्रबंधन करता है क्योंकि ML मॉडल निरंतर अनुमान लगाते हैं। संकलन-समय और रनटाइम सॉफ़्टवेयर दोनों के लिए, थीसिस उद्योग और शिक्षा जगत से मौजूदा सॉफ़्टवेयर स्टैक का सावधानीपूर्वक ऑडिट करती है, उनकी सीमाओं की पहचान करती है। इस विश्लेषण के आधार पर, इन कमियों को दूर करने और समग्र सिस्टम प्रदर्शन को बढ़ाने के लिए नए सॉफ़्टवेयर समाधान डिज़ाइन और कार्यान्वित किए जाते हैं। इसके अतिरिक्त, थीसिस Nvidia एज GPU के लिए सिस्टम सॉफ़्टवेयर के विभिन्न पहलुओं की खोज करती है, जिसमें एप्लिकेशन की गतिशीलता और परिवेश के तापमान को प्रबंधित करने की तकनीकें शामिल हैं, और कई विक्रेताओं से एप्लिकेशन तैनात करने की चुनौतियों का समाधान करती है।

हम सबसे पहले एज इंफ़रेंस के लिए पूर्व-प्रशिक्षित तंत्रिका नेटवर्क को संपीड़ित करने के लिए प्रमुख सॉफ़्टवेयर फ़्रेमवर्क की जाँच करते हैं, जिसमें TensorRT, AutoTVM और AutoScheduler पर ध्यान केंद्रित किया जाता है। हम इन फ़्रेमवर्क की तुलना करने के लिए मेट्रिक्स का एक व्यापक सेट विकसित करते हैं और उन मेट्रिक्स को अनुभवजन्य रूप से मापने के लिए सावधानीपूर्वक प्रयोग डिज़ाइन करते हैं। TensorRT के हमारे विश्लेषण ने उल्लेखनीय प्रदर्शन और सटीकता में सुधार दिखाया, लेकिन अधिक शक्तिशाली हार्डवेयर पर अलग-अलग आउटपुट और बढ़ी हुई विलंबता जैसे अप्रत्याशित गैर-निर्धारक व्यवहार भी प्रकट किए। हमने विक्रेता-अज्ञेय फ़्रेमवर्क AutoTVM और AutoScheduler का भी मूल्यांकन किया, उनकी तुलना प्रदर्शन की भविष्यवाणी, अनुमान विलंबता, आउटपुट पुनरुत्पादन और हार्डवेयर उपयोग पर की। इस तुलना से कई दिलचस्प अवलोकन हुए, जिनका ट्रैफ़िक मॉनिटरिंग जैसे शहरी स्थिरता अनुप्रयोगों के लिए महत्वपूर्ण निहितार्थ हैं। ये जानकारीयां वास्तविक समय,

सुरक्षा-महत्वपूर्ण अनुप्रयोगों के लिए सबसे प्रभावी DNN कंपाइलर का चयन करने के लिए सिफारिशें प्रदान करती हैं, जिससे सटीक और पूर्वानुमानित निष्कर्ष सुनिश्चित होते हैं।

इसके अलावा, विकासशील क्षेत्रों की चरम स्थितियों में उच्च थ्रूपुट को प्रबंधित करने और सिस्टम की दीर्घायु सुनिश्चित करने की चुनौतियों का समाधान करने के लिए, हमने DynCNN विकसित किया, जो अनुप्रयोग गतिशीलता और परिवेश तापमान प्रबंधन के लिए डिज़ाइन किया गया एक नियंत्रक है। DynCNN थ्रेड काउंट और एक्सेलेरेटर आवृत्तियों को विनियमित करने के लिए प्रोसेसर विषमता का लाभ उठाता है, सख्त थर्मल और पावर बाधाओं का पालन करते हुए अनुप्रयोग प्रदर्शन को अनुकूलित करता है। हमने वास्तविक ट्रैफिक इंटरसेक्शन से 40-दिवसीय डेटासेट का उपयोग करके तीन व्यावसायिक रूप से उपलब्ध एम्बेडेड GPU पर DynCNN का मूल्यांकन किया, जो इन मांग वाली स्थितियों के तहत सिस्टम दक्षता बढ़ाने में इसकी प्रभावशीलता को प्रदर्शित करता है।

इसके अलावा, एज डिवाइस को कई विक्रेताओं से एप्लिकेशन तैनात करने में चुनौतियों का सामना करना पड़ता है, जिनमें से प्रत्येक की अपनी सॉफ्टवेयर निर्भरताएँ होती हैं। Docker कंटेनर का उपयोग अलगाव प्रदान करता है और इन निर्भरताओं को प्रबंधित करता है, जिससे सुरक्षा और पोर्टेबिलिटी बढ़ती है। हालाँकि, एक साथ चलने वाले एप्लिकेशन एक-दूसरे के प्रदर्शन को प्रभावित कर सकते हैं और साझा संसाधन विवाद के कारण डिवाइस का तापमान बढ़ा सकते हैं। इसके लिए इष्टतम प्रदर्शन और थर्मल प्रबंधन सुनिश्चित करने के लिए उन्नत शेड्यूलिंग की आवश्यकता होती है। इन मुद्दों को संबोधित करने के लिए, हम CPU-GPU आधारित एज कंप्यूटिंग सर्वर के लिए एक अनुमान सेवारत फ्रेमवर्क, IceEdge का प्रस्ताव करते हैं। आइसएज आइसोलेशन के लिए कंटेनर का उपयोग करता है और इसमें एमएल वर्कलोड के लिए डेडलाइन- और थर्मल-अवेयर शेड्यूलर शामिल है, जो अच्छे GPU उपयोग को सुनिश्चित करता है और सेवा की गुणवत्ता को बनाए रखता है। आइसएज में एक नियंत्रक भी शामिल है जो CPU-GPU संचालन को सटीक रूप से प्रबंधित करता है, OS थर्मल नियंत्रण पर निर्भरता से बचता है और इस प्रकार दक्षता और स्थिरता में सुधार करता है। गतिशील कार्यभार और एज प्लेटफॉर्म की थर्मल सुरक्षा बाधाओं दोनों पर विचार करते हुए क्लॉक फ्रीक्वेंसी।

निष्कर्ष में, यह थीसिस विकासशील देशों में स्थिरता अनुप्रयोगों के लिए Nvidia एज GPU प्लेटफॉर्म पर डीप न्यूरल नेटवर्क (DNN) की तैनाती की जाँच करती है। हमारा उद्देश्य एक व्यापक समाधान प्रदान करना है जो संकलक चयन, थर्मल प्रबंधन और बहु-विक्रेता अनुप्रयोग परिनियोजन जैसी महत्वपूर्ण चुनौतियों का समाधान करके एज कंप्यूटिंग सिस्टम की दीर्घायु, प्रदर्शन और विश्वसनीयता को बढ़ाता है। हमारे शोध से प्राप्त अंतर्दृष्टि और सिफारिशें मजबूत और कुशल एज कंप्यूटिंग समाधानों के निर्माण, सतत विकास को बढ़ावा देने और अविकसित क्षेत्रों में जीवन की गुणवत्ता में सुधार करने में सहायता करेंगी।

Contents

Certificate	i
Acknowledgements	iii
Abstract	vii
LIST OF FIGURES	xxiii
LIST OF TABLES	xxviii
1 INTRODUCTION	1
1.1 Promise Of Edge Computing	3
1.2 Challenges in Urban Traffic Management for Nvidia Edge GPU Platforms . . .	4
1.3 Contribution of the Thesis	6
1.3.1 A Quantitative Audit of NN Compilers	6
1.3.2 Managing Application Dynamism and Thermal Management	8
1.3.3 Managing Multi-Vendor Application Deployment and Resource Man- agement	9
1.4 Organization of the Thesis	10
2 A QUANTITATIVE AUDIT OF NN COMPILERS	11

2.1	Evaluation Platforms Used	12
2.2	Neural Network Compilers	13
2.3	Demystifying TensorRT: Characterizing Neural Network Inference Engine on Nvidia Edge Devices	16
2.3.1	Measurement Setup	16
2.3.2	Evaluated Neural Network (NN) Models	16
2.3.3	Application Profiling Softwares	17
2.3.4	Image Datasets and NN Model Outputs	17
2.3.5	Evaluation Metrics	17
2.3.6	Experimental Methodology	18
2.3.7	Summary of Key Findings on TensorRT engines	19
2.3.8	Accuracy Metric Analysis	19
2.3.9	Classification Accuracy on Benign and Adversarial Datasets	20
2.3.10	Consistency of Output Labels	21
2.3.11	Performance Metric Analysis	22
2.3.12	Throughput for Image Classification Networks	23
2.3.13	Throughput for Object Detection Models and Concurrency	23
2.3.14	Inference Latency for Classification and Detection Networks	26
2.3.15	Inference Latency Anomaly Analysis	28
2.3.16	Latency Anomalies of same TensorRT Engines across Platforms	28
2.3.17	Latency Differences Across TensorRT engines on Same Platform	28
2.3.18	Implications of Findings	30
2.3.19	Impact on Applications	30

2.3.20	Impact on Micro-architecture Performance Modeling	32
2.3.21	Related Work	36
2.3.22	Conclusion and Future Work	36
2.4	Implications of Utilizing DNN Compilers on Edge GPUs for Real-Time and Safety-Critical Systems	37
2.4.1	Experimental Setup	37
2.4.2	Are Inferences Accurate and Reproducible?	38
2.4.3	Do optimizations retain the pre-trained model accuracy?	38
2.4.4	Do optimized engines give consistent output on same input?	39
2.4.5	Are Inference Latencies Low and Predictable?	42
2.4.6	Comparison of Inference Latencies across NN compilers	43
2.4.7	Inference Latency Predictability Across Compiled Engines	44
2.4.8	Effect of Dynamic Voltage and Frequency Scaling (DVFS) on Inference Latency	47
2.4.9	Comparison of Compilation Latency across NN compilers	49
2.4.10	Is Resource Utilization Optimal for NN Compilers?	50
2.4.11	Impact Analysis on Real Time Applications	52
2.4.12	Safety Critical Real Time Application Instances	52
2.4.13	Response Time Components for Deployed DNN Models	53
2.4.14	Response Times vs. NN Compilers For Single Camera	53
2.4.15	Response Times vs. NN Compilers For Multiple Cameras	53
2.4.16	Worst Case Response or Execution Time (WCET) Analysis	56
2.4.17	Related Work	57
2.4.18	Recommendations	59

2.4.19	Conclusion	59
3	APPLICATION DYNAMISM AND AMBIENT TEMPERATURE AWARE NEURAL NETWORK SCHEDULER IN EDGE DEVICES	61
3.1	The Promise of Concurrent Neural Networks on Embedded GPU Devices . . .	62
3.1.1	Neural Network Concurrency with Multi-processing	62
3.1.2	Neural Network Concurrency with Multi-threading	63
3.2	Developing Region Traffic Monitoring Case-study	65
3.2.1	Concurrent CNNs for traffic monitoring and rule violation detection . .	65
3.2.2	Traffic density estimation on embedded CPU cores	67
3.2.3	Long term traffic density monitoring to understand application dynamism	69
3.3	Thermal and Power Constraints in Neural Network Concurrency	70
3.3.1	Platform temperature and power consumption rises as GPU thread count increases	71
3.3.2	Platform temperature and power consumption rises as GPU frequency increases	71
3.3.3	Platform temperature rises more as ambient temperature rises	72
3.4	DynCNN: Application dynamism aware CNN concurrency controller	74
3.4.1	Can variable traffic densities measured by CPU control GPU concurrency?	75
3.4.2	DynCNN Thermal Control Algorithm	77
3.4.3	DynCNN Controller Evaluation	78
3.5	Related Work	84
3.6	Discussion	86
3.7	Conclusion and Future Work	86

4	ICEEDGE: THERMAL-AWARE MACHINE LEARNING INFERENCE SERV- ING FOR SUPPORTING MULTI-TENANCY	89
4.1	IceEdge Design Choices	90
4.2	IceEdge System Design	93
4.2.1	Scheduling	94
4.2.2	Thermal Control	97
4.3	Implementation	99
4.4	Evaluation	100
4.4.1	Bare-Metal vs. Docker	101
4.4.2	IceEdge Latency Improvement	102
4.4.3	Profiling Table Updates in IceEdge	104
4.4.4	IceEdge vs. Baselines: Thermal Safety	106
4.4.5	Traffic Monitoring Case Study	106
4.5	Related Work	112
4.6	Conclusion and Future Work	113
5	FUTURE WORK AND CONCLUSION	115
5.1	Future Work	115
5.2	Conclusion	117
	Bibliography	119
	List of Publications	133
	Biography	135

List of Figures

1.1	Thesis Outline	6
2.1	Hardware and software eco-system for DNN execution in real time applications like ADAS [1, 2, 3] and industrial control [4].	11
2.2	TensorRT, AutoTVM and AutoScheduler optimizations	14
2.3	FPS and GPU utilizations on NX and AGX when Tiny-Yolov3 CNN is run. NX saturates at 28 GPU threads and AGX at 36 GPU threads, with GPU utilization slightly above 80% in both cases. RAM bandwidth bottleneck marks this thread saturation points. FPS is around 196 and 346 per thread for the saturation thread count, at 1109.25 MHz GPU frequency on NX and 1377 MHz on AGX.	24
2.4	FPS and GPU utilizations on NX and AGX when Googlenet CNN is run. NX at 16 GPU threads and AGX at 24 GPU threads, with GPU utilization slightly between 82-85% for the two boards. RAM bandwidth bottleneck marks this thread saturation points. FPS is around 85 and 210 per thread for the saturation thread count, at 1109.25 MHz GPU frequency on NX and 1377 MHz on AGX.	25
2.5	Small probability differences for output predictions for engine1 and engine2 in Resnet-18.	41
2.6	GPU utilization for ResNet-18 and GoogleNet for DVFS enabled and fixed frequency on AGX	48
2.7	Energy utilization (joules/image) for ResNet-18 and GoogleNet for DVFS enabled and fixed frequency on AGX	48
2.8	Inference time(ms) for ResNet-18 on AGX	54

2.9	Response time components across NN compilers for 500 ImageNet samples (Resnet-18 on AGX)	54
2.10	Image read, pre-processing and post-processing times are similar for different NN compilers for 500 different images of ImageNet dataset, with increasing camera feeds. This is representative execution of Resnet-18 on AGX platform. .	55
2.11	Inference time(in ms) in multiple camera feeds for image classification and object detection networks used in ADAS	55
3.1	Number of concurrent processes is limited by RAM size. The maximum processes that are supported on TX2 are 8.	63
3.2	Different GPU concurrency techniques where multi-processing sees reduction in FPS due to context-switching overhead.	63
3.3	FPS and GPU utilizations on TX2, NX and AGX when Tiny-YOLO CNN is run. TX2 saturates at 24 GPU threads, NX at 28 GPU threads and AGX at 36 GPU threads, with GPU utilization slightly above 80% in all three cases. RAM bandwidth bottleneck marks this thread saturation points. FPS is around 64, 196 and 346 per thread for the saturation thread count, depending on the supported clock frequencies of the three platforms.	64
3.4	FPS and GPU utilizations on TX2, NX and AGX when Googlenet CNN is run. TX2 saturates at 14 GPU threads, NX at 16 GPU threads and AGX at 24 GPU threads, with GPU utilization slightly between 82-85% for the three boards. RAM bandwidth bottleneck marks this thread saturation points. FPS is around 55, 85 and 210 per thread for the saturation thread count, depending on the supported clock frequencies of the three platforms.	64
3.5	A Typical 4-way intersection	66
3.6	View from traffic cameras on the left and traffic densities measured on road deployed CPU on the right	68
3.7	Box plot showing traffic queue density for 40 days. Traffic density increases or decreases at different hours on different days. Road approach 1 and 2 have higher traffic densities compared to approach 3 where traffic gradually increases and then continues to remain high.	69

3.8	Temperature and Power with varying CNN thread counts for AGX (Representative graph). Temperature rises upto 20 degree Celsius and Power varies by 7500 mW on AGX, as GPU thread count changes from 1 to 24.	71
3.9	Temperature and Power with varying GPU frequencies for AGX (Representative graph). Temperature rises upto 25 degree Celsius and Power varies by 10000 mW on AGX, as GPU frequency changes from the minimum to the maximum.	72
3.10	Temperature rise at 35 degree ambient temperature. Keeping the fan on throughout maintains the temperature lower by around 10 degrees.	73
3.11	Temperature rise at 45 degree ambient temperature. Keeping the fan on throughout maintains the temperature lower by around 10 degrees.	73
3.12	DynCNN Thermal Controller Workflow.	77
3.13	Application performance in FPS with SchedUtil CPU governor as CPU Baseline1 and multiple GPU DVFS governors given in Table 3.3. All the governors except OnDemand maintain the higher application FPS (more by 10 in TX2, 50 in NX and AGX) than DynCNN. OnDemand maintains itself at a lower frequency due to battery savings and hence has lower FPS	80
3.14	Average temperature with SchedUtil CPU governor as CPU Baseline1 and multiple GPU DVFS governors given in Table 3.3. Average CPU and GPU temperatures for DynCNN are less by 5-6°C for TX2, 2-10 °C for NX and 2-22 °C for AGX as all the existing governors target the application performance by keeping the frequency maximum most of the times.	81
3.15	Average power with SchedUtil CPU governor as CPU Baseline1 and multiple GPU DVFS governors given in Table 3.3. The average power consumed by DynCNN is less by 4000 mW,1300-5000 mW and 5000-14500 mW for TX2, NX and AGX respectively. DynCNN dynamically adjusts the frequency and the count of threads leading to less power consumption.	81
3.16	Application performance in FPS with fixed CPU frequencies as CPU Baseline2 and multiple GPU DVFS governors given in Table 3.3. Existing DVFS governors adjust the GPU frequency based on the setting of CPU frequency. Performance improvement of 31.2% is observed compared to the existing governors.	83

3.17	Average temperature with fixed CPU frequencies as CPU Baseline2 and multiple GPU DVFS governors given in Table 3.3. All the existing governors maintain the GPU at a lower frequency due to setting of CPU at a fixed frequency and hence average temperature is less by 2-4°C compared to DynCNN.	83
3.18	Average power with fixed CPU frequencies as CPU Baseline2 and multiple GPU DVFS governors given in Table 3.3. GPU frequency maintained is lower in existing governors and hence power consumed is less compared to DynCNN.	83
4.1	(a) inference latency (c) GPU frequency vs. temperature (d) throughput and GPU utilisation	92
4.2	IceEdge design and workflow. 1 Sensor inputs are sent to a model in a Docker container. 2 The container forwards request metadata to a central controller. 3 The controller determines the schedule for requests and sends the schedule to respective containers. 4 The container forwards the requests to the GPU for inference as per schedule.	93
4.3	Frame drop percentage in different frameworks for varying number of co-located models.	103
4.4	Frequency count of the batch configurations for two models (Resnet-18 and Inception) when run concurrently.	104
4.5	Comparison of change in platform temperatures when running Triton, DeepRT, and IceEdge. IceEdge explicitly sets the GPU frequencies and switches based on the queue density and the rate of change of temperature.	107
4.6	(a) Frame drop percentage (b) 2 min trace from six cameras (c) Inference request arrival pattern	108
4.7	(a) batch sizes for TinyYolo (b) batch sizes for Resnet-18 (c) Concurrency in IceEdge	109
4.8	Temperature rise when processing data with varying traffic densities averaged over 40 days. IceEdge maintains the temperature well below the thermal trip point (80°C)	110

4.9	Frame drop percentage comparison when we scale the number of cameras from 8 to 32. Each camera needs 30FPS as we assume the empty roads for this experiment	111
-----	---	-----

List of Tables

2.1	Evaluation platforms with Embedded NVIDIA GPU	13
2.2	Neural network (NN) models used for our evaluation, with model sizes (MB) with and without TensorRT optimizations.	16
2.3	Summary of empirical findings on TensorRT engines	19
2.4	Top-1 Error(%) for image classification networks on benign dataset using Ten- sorRT optimized and un-optimized engines.	20
2.5	Top-1 Error(%) for image classification networks on adversarial dataset using TensorRT optimized and un-optimized engines.	20
2.6	Number of different prediction output across engines for different platforms (cross platform engines).	21
2.7	Number of different prediction output across engines (platform specific en- gines).	22
2.8	FPS for TensorRT optimized and un-optimized engines	23

2.9	Average run time or inference latency (in ms) of different networks along with standard deviation, when run with TensorRT optimizations. Platform AGX with more hardware resources and slightly higher GPU clock frequency during experiments, should have lower inference latencies than platform NX. When this expectation does not hold, an anomaly occurs. There are 3 possible anomalous cases of higher AGX runtime (rAGX) than NX runtime (rNX) - case ❶ (highlighted in bold font) when two different engines are compiled on NX and AGX and each platform runs its own compiled engine, case ❷ (highlighted in blue color) when both platforms run the same engine compiled on NX and case ❸ (highlighted in red color) when both platforms run the same engine compiled on AGX.	24
2.10	Average inference time (in ms) along with standard deviation, when TensorRT engines are run without nvprof tool.	27
2.11	Average run time (in ms) with CUDA mempcy time included and excluded along with standard deviation	27
2.12	Average run time (in ms) of similar set of kernels on NX and AGX along with standard deviation	27
2.13	Average run time (in ms) using different TensorRT engines of the same NN models for AGX platform. Blue rows indicate run time differences across the three engines.	29
2.14	Number of invocations and run time (in μ s) per invocation of a CUDA kernel in <i>inception-v4</i> on AGX platform.	30
2.15	TensorRT positive impact on automotive applications	31
2.16	TensorRT negative impact on automotive applications	32
2.17	Different λ values obtained across 3 engines of the <i>inception-v4</i> for different kernels along with the associated prediction error in percentage.	34
2.18	Prediction Error(%) for some kernels of <i>Mobilenetv1</i> along with two different λ s (one obtained using metrics of NX and other using metrics of AGX)	35
2.19	Neural network models used in the study	38

2.20	Top-1 Accuracy Comparison (in %) across different NN compiled models and unoptimized NN models.	39
2.21	Top-1 mismatches across TensorRT engines	40
2.22	Top-5 consistency across TensorRT engines	40
2.23	Output mismatches for FP32 TensorRT engines	40
2.24	Output mismatches across inference engines in TensorRT using caching technique	41
2.25	Inference Latency (ms) across NN compilers	43
2.26	Inference times (ms) for Auto-Scheduler engines on V100, with increasing search iterations.	44
2.27	Inference latency (ms) across AutoTVM engines	44
2.28	Inference time (ms) across TensorRT engines on AGX	44
2.29	Inference time (ms) across AutoScheduler engines on AGX	45
2.30	Layer to Kernel mapping along with measurement time (μ s) for each kernel during engine building for TensorRT. Engine1 (highlighted in blue) and engine2 (highlighted in red) pick two different kernels for the same layer.	46
2.31	Inference time (ms) across TensorRT engines on AGX using caching	46
2.32	Sample kernels with execution times and % of total execution time spent in them, for two AutoScheduler engines.	46
2.33	Inference Latency (ms) across NN compilers when DVFS is disabled on AGX	47
2.34	Engine building time (secs) on AGX	49
2.35	Inference latencies (in ms) for NX vs AGX	50
2.36	Response latencies (ms) across NN compilers	55
2.37	Desirable behaviors of NN compilers and their impact on intersection control and ADAS applications. ✓represents that the NN compiler shows the desired behavior, ✗indicates it doesn't. As TensorRT has more ✗, AutoTVM/AutoScheduler can be recommended for real time application settings.	58

3.1	6 CNN tasks per approach at an intersection, adding up to 18 or 24 CNN tasks for a 3-way or 4-way intersection	67
3.2	Traffic density estimation at an intersection that can execute on CPU	68
3.3	Existing GPU governors used as Baselines to evaluate DynCNN metric trade-offs	79
4.1	DNNs used in our evaluation study and their avg inference time (ms) for a batch size of 4.	102
4.2	Profiling table updates of execution times (in ms) for IceEdge controller after every 2000 iterations.	105