

THE UNIVERSITY OF TULSA
THE GRADUATE SCHOOL

GENERATION TECHNIQUES AND MESSAGE-PASSING INTERFACE
IMPLEMENTATION FOR COMPLIANCE GRAPHS

by
Noah L. Schrick

A thesis submitted in partial fulfillment of
the requirements for the degree of Master of Science
in the Discipline of Computer Science

The Graduate School
The University of Tulsa

2022

THE UNIVERSITY OF TULSA
THE GRADUATE SCHOOL

GENERATION TECHNIQUES AND MESSAGE-PASSING INTERFACE
IMPLEMENTATION FOR COMPLIANCE GRAPHS

by
Noah L. Schrick

A THESIS
APPROVED FOR THE DISCIPLINE OF
COMPUTER SCIENCE

By Thesis Committee

Peter J. Hawrylak, Chair
John Hale
Mauricio Papa

COPYRIGHT STATEMENT

Copyright © 2022 by Noah L. Schrick

All rights reserved. No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, recording, or otherwise, without the prior written permission of the author.

ABSTRACT

Noah L. Schrick (Master of Science in Computer Science)

Generation Techniques and Message-Passing Interface Implementation for Compliance Graphs

Directed by Peter J. Hawrylak

73 pp., Chapter 6: Conclusions and Future Works

(196 words)

Attack graphs have historically been used to represent the state of a system or set of systems, and illustrate ways that attackers can carry out exploits to put the systems in a critical position. By refining the information and format of attack graphs, a similar process can be applied for compliance and regulation representations to illustrate ways that systems may violate or fall out of compliance in a format called compliance graphs. Like attack graphs, compliance graphs also suffer from the state space explosion problem, and generated graphs quickly become very large even for small networks. This work introduces extensions to an attack graph generator, RAGE attack graph engine (RAGE), to support compliance graph generation, and these extensions have led to successfully generating compliance graphs that can then be analyzed through an independent process. Additional extensions have been introduced to support the utility of RAGE, and this includes the implementation of a synchronous firing feature to prevent generation of states where assets deviate from a shared, inseparable feature such as time. The compliance graph generation algorithm has also been modified by two different approaches to expand the generation process to function for distributed computing environments using Message Passing Interface (MPI).

ACKNOWLEDGEMENTS

I would like to acknowledge and give my sincerest gratitude to my wife for all of her continued support and patience through this process. I would also like to extend this appreciation to all my family and friends, who have provided constant encouragement and motivation during my work. I would also like to greatly thank my advisor, Dr. Hawrylak, for his guidance, support, and motivation during my time through this program. He has served as a remarkably vital role not only through this thesis, but also for the advice and encouragement offered throughout my graduate career. Further thanks to the members of my committee, Dr. Papa and Dr. Hale, who have also provided their expertise and support throughout this journey.

TABLE OF CONTENTS

COPYRIGHT	iii
ABSTRACT	iv
ACKNOWLEDGEMENTS	v
TABLE OF CONTENTS	viii
LIST OF TABLES	ix
LIST OF FIGURES	xi
CHAPTER 1: INTRODUCTION	1
1.1 Introduction to Attack Graphs	1
1.2 Application to Compliance	2
1.2.1 <i>Introduction to Compliance Graphs</i>	2
1.2.2 <i>Defining Compliance Graphs</i>	3
1.2.3 <i>Difficulties of Compliance Graphs</i>	3
1.3 Objectives and Contributions	4
CHAPTER 2: RELATED WORKS	5
2.1 Introduction to Graph Generation	5
2.2 Graph Generation Improvements	5
2.3 Improvements Specific to Attack Graph Generation	6
CHAPTER 3: UTILITY EXTENSIONS TO THE RAGE ATTACK GRAPH GENERATOR	8
3.1 Path Walking	8
3.2 Color Coding	9
3.3 Compound Operators	10
3.4 Relational Operators	12
3.5 Intermediate Database Storage	13
3.5.1 <i>Introduction to Intermediate Database Storage</i>	13
3.5.2 <i>Memory Constraint Difficulties</i>	13
3.5.3 <i>Maximizing Performance with Intermediate Database Storage</i>	15
3.5.4 <i>Portability</i>	16

CHAPTER 4: SYNCHRONOUS FIRING	18
4.1 Introduction	18
4.1.1 <i>Related Synchronous Firing Work</i>	19
4.2 Necessary Alterations and Additions	20
4.2.1 <i>GNU Bison and Flex</i>	20
4.2.2 <i>PostgreSQL</i>	21
4.2.3 <i>Compound Operators</i>	22
4.2.4 <i>Graph Generation</i>	22
4.3 Experimental Networks and Results	23
4.3.1 <i>Experimental Networks</i>	23
4.3.2 <i>Results</i>	25
CHAPTER 5: PARALLELIZATION USING MESSAGE PASSING INTER- FACE	29
5.1 Introduction to MPI Utilization for Attack and Compliance Graph Generation	29
5.2 Necessary Components	29
5.2.1 <i>Serialization</i>	29
5.3 Tasking Approach	30
5.3.1 <i>Introduction to the Tasking Approach</i>	30
5.3.2 <i>Algorithm Design</i>	32
5.3.2.1 Communication Structure	34
5.3.2.2 Task 0	34
5.3.2.3 Task 1	34
5.3.2.4 Task 2	35
5.3.2.5 Task 3	37
5.3.2.6 Task 4 and Task 5	37
5.3.2.7 MPI Tags	38
5.3.3 <i>Performance Expectations and Use Cases</i>	38
5.3.4 <i>Results</i>	39
5.4 Subgraphing Approach	40
5.4.1 <i>Introduction to the Subgraphing Approach</i>	42
5.4.2 <i>Algorithm Design</i>	42
5.4.2.1 Worker Nodes	43
5.4.2.2 Root Node	44
5.4.2.3 Database Node	45
5.4.2.4 MPI Tags	45
5.4.3 <i>Performance Expectations and Use Cases</i>	45
5.4.4 <i>Results</i>	46
CHAPTER 6: CONCLUSIONS AND FUTURE WORKS	55
6.1 Conclusions	55
6.2 Future Work	56
NOMENCLATURE	59

BIBLIOGRAPHY	59
------------------------	----

LIST OF TABLES

5.1	MPI Tags for the MPI Tasking Approach	39
5.2	MPI Tags for the MPI Subgraphing Approach	46

LIST OF FIGURES

3.1	Path Walking to State 14	9
3.2	Color Coding a Small Network Based on Violations	10
4.1	A network without Synchronous Firing generating infeasible states	19
4.2	Inclusion of Synchronous Firing into GNU Bison, GNU Flex, and the overall program	22
4.3	Synchronous Firing in the Graph Generation Process	24
4.4	Bar Graph and Line Graph Representations of Synchronous Firing on Runtime	27
4.5	Bar Graph and Line Graph Representations of Synchronous Firing on State Space	28
5.1	Task Overview of the Attack and Compliance Graph Generation Process . . .	31
5.2	Node Allocation for each Task	33
5.3	Data Distribution of Task One	35
5.4	Communication From Task 1 to Task 2 when the Number of Nodes Allocated is Equal	36
5.5	Communication From Task 1 to Task 2 when Task 1 Has More Nodes Allocated	37
5.6	Example of a Not Applicable Exploit for the MPI Tasking Testing	41
5.7	Speedup and Efficiency of the MPI Tasking Approach for a Varying Number of Compute Nodes with an Increasing Problem Size	41
5.8	Example Graph Using the MPI Subgraphing Approach	43
5.9	Frontier Merging and Data Distribution Process	45
5.10	First iteration results of MPI Subgraphing in terms of Runtime	47
5.11	First iteration results of MPI Subgraphing in terms of Speedup and Efficiency	48

5.12	Modified Subgraphing Example Graph with Two New Edges	50
5.13	Duplicate States Explored vs Actual Number of States for the 1-4 Service Tests	51
5.14	Speedup and Efficiency of MPI Subgraphing when using a DHT	53
5.15	Runtime of MPI Subgraphing when using a DHT vs not using a DHT	54
6.1	Possible Method for Blending MPI and OpenMP for Task 2 of the MPI Tasking Approach	57

CHAPTER 1

INTRODUCTION

1.1 Introduction to Attack Graphs

Cybersecurity has been at the forefront of computing for decades, and vulnerability analysis modeling has been utilized to mitigate threats to aid in this effort. One such modeling approach is to represent a system or a set of systems through graphical means, and encode information into the nodes and edges of the graph. Even as early as the late 1990s, experts have composed various graphical models to map devices and vulnerabilities through attack trees, and this work can be seen through the works published by the authors of [1]. This work, and other attack tree discussions of this time such as that conducted by the author of [2], would later be referred to as early versions of modern-day attack graphs [3]. By utilizing this graphical approach, cybersecurity postures can be measured at a system's current status, as well as hypothesize and examine other postures based on system changes over time.

Attack graphs are an appealing approach since they are often designed to be exhaustive: all system properties are represented at its initial state, all attack options are fully enumerated, all permutations are examined, and all changes to a system are encoded into their own independent states, where these states are then individually analyzed through the process. The authors of [4] also discuss the advantage of conciseness of attack graphs, where the final graph only incorporates states that an attacker can leverage; no superfluous states are generated that can clutter analysis. Despite their advantages, attack graphs do suffer from their exhaustiveness as well. As the authors of [3] examine, even very small networks with only 10 hosts and 5 vulnerabilities yield graphs with 10 million edges. When scaling attack graphs to analyze the modern, interconnected state of large networks comprising of a

multitude of hosts, and utilizing the entries located in the National Vulnerability Database and any custom vulnerability testing, attack graph generation quickly becomes infeasible. Similar difficulties arise in related fields, where social networks, bioinformatics, and neural network representations result in graphs with millions of states [5]. Various efforts that will be discussed in Section 2.3 demonstrate methods and techniques that can mitigate these difficulties and improve performance.

1.2 Application to Compliance

1.2.1 Introduction to Compliance Graphs

As an alternative to attack graphs for examining vulnerable states and measuring cybersecurity postures, the focus can be narrowed to generate graphs with the purpose of examining compliance or regulation statuses. These graphs are known as compliance graphs. Compliance graphs can be especially useful for cyber-physical systems, where a greater need for compliance exists. As the authors of [6], [7], and [8] discuss, cyber-physical systems have seen greater usage, especially in areas such as critical infrastructure and Internet of Things. The challenge of cyber-physical systems lies not only in the demand for cybersecurity of these systems, but also the concern for safe, stable, and undamaged equipment. The industry in which these devices are used can lead to additional compliance guidelines that must be followed, increasing the complexity required for examining compliance statuses. Compliance graphs are promising tools that can aid in minimizing the overhead caused by these systems and the regulations they must follow.

A few alterations are required for attack graph generators to function as compliance graph generators, and these alterations are discussed in Section 1.2.2. Compliance requirements are broad and varying, and can function as safety regulations, maintenance compliance, or any other regulatory compliance. In the same fashion as attack graphs, compliance graphs are exhaustive, and future system states can be analyzed to determine appropriate steps that need to be taken for preventative measures [6].

1.2.2 *Defining Compliance Graphs*

The common encodings of attack graph properties have different encoding definitions in compliance graphs. The nodes of an attack graph typically represent the system state that includes the qualities and topologies of all assets in the network as they pertain to cybersecurity postures. Nodes of a compliance graphs also represent the system state, however they include the qualities and topologies of all assets in the network as they pertain to compliance regulation. For instance, a quality for a vehicle’s maintenance compliance could be described as: *car:months_since_oil_change=6*, or *car:miles_since_oil_change=10,000*. Edges represent changes to a system state that inserted, modified, or deleted a quality or topology. Using the car example, an edge could represent the addition of more mileage or more time since the last oil change. One large differentiation of attack graphs and compliance graphs can be seen through topologies. For assets in attack graphs, topologies typically represent a connection of assets through a digital medium. For compliance graphs, topologies not only need to represent the digital connections of assets, but also need extensions to incorporate hardware devices such as sensors, actuators, or other physical equipment [6]. In addition, rather than using applicable exploits or vulnerabilities to an asset, compliance violation detections should be used. An attack graph generation engine would need to use compliance parameters rather than exploit files, but would otherwise function similarly in the generation process.

1.2.3 *Difficulties of Compliance Graphs*

Like attack graphs, compliance graphs suffer from the state space explosion problem. Since compliance graphs are also exhaustive, the resulting networks can grow to incredibly large sizes. Compliance regulations that need to be checked at each system state such as SOX [9], HIPAA [10], GDPR [11], PCI DSS [12], or any other regulatory compliance in conjunction with a large number of assets that need to be checked can very quickly produce large resulting graphs. The creation of these graphs through a serial approach likewise becomes increasingly infeasible in a reasonable amount of time. Due to this, the

high-performance computing (HPC) space presents itself as an appealing approach. This work aims to extend the attack graph generator engine RAGE in [13] to begin development for compliance graph generation. The example networks in this work will also be in the compliance graph space, specifically examining vehicle maintenance compliance. This work will also examine approaches to leverage high-performance computing to aid in the generation of compliance graphs.

1.3 Objectives and Contributions

The objectives of this thesis are:

- Extend the utility of RAGE to:
 1. Reduce the complexity required for network model and exploit file creation
 2. Expand the complexity of attack modeling
 3. Allow for the creation of a very large sized attack graph, assuming very large storage
 4. Split Attack Graphs into subgraphs to simplify analysis of individual clusters
- Implement a synchronous firing solution to reduce state space explosion for inseparable features while remaining exhaustive and capturing all necessary information
- Extend RAGE to function for distributed computing environments
- Extend RAGE to support compliance graph generation

CHAPTER 2

RELATED WORKS

Many researchers have developed or extended attack graphs since their beginning as attack trees. This Chapter reviews a few of their efforts as they relate to this work and to graph generation.

2.1 Introduction to Graph Generation

Graph generation as a broad topic has many challenges that prevent full actualization of computation seen from a theoretical standpoint. In actuality, graph generation often achieves only a very low percentage of its expected performance [14]. A few reasons for this occurrence lie in the underlying mechanisms of graph generation. The generation is predominantly memory based (as opposed to based on processor speed), where performance is tied to memory access time, the complexity of data dependency, and coarseness of parallelism [14], [5], [15]. Graphs consume large amounts of memory through their nodes and edges, graph data structures suffer from poor cache locality, and memory latency from the processor-memory gap all slow the generation process dramatically [14], [15]. Section 2.2 discusses a few works that can be used to improve the graph generation process, and Section 2.3 discusses a few works specific to attack graph generation improvements.

2.2 Graph Generation Improvements

For architectural and hardware techniques for generation improvement, the authors of [15] discuss the high cache miss rate, and how general prefetching does not increase the prediction rate due to nonsequential graph structures and data-dependent access patterns. However, the authors continue to discuss that generation algorithms are known in advance, so explicit tuning of the hardware prefetcher to follow the traversal order pattern can lead to

better performance. The authors were able to achieve over 2x performance improvement of a breadth-first search approach with this method. Another hardware approach is to make use of accelerators. The authors of [16] present an approach for minimizing the slowdown caused by the underlying graph atomic functions. By using the atomic function patterns, the authors utilized pipeline stages where vertex updates can be processed in parallel dynamically. Other works, such as those by the authors of [5] and [17], leverage field-programmable gate arrays (FPGAs) for graph generation in the HPC space through various means. This includes reducing memory strain, storing repeatedly accessed lists, storing results, or other storage through the on-chip block RAM, or even leveraging Hybrid Memory Cubes for optimizing parallel access.

From a data structure standpoint, the authors of [18] describe the infeasibility of adjacency matrices in large-scale graphs, and this work and other works such as those by the authors of [19] and [20] discuss the appeal of distributing a graph representation across systems. The author of [20] discusses the usage of distributed adjacency lists for assigning vertices to workers. The authors of [20] and [21] present other techniques for minimizing communication costs by achieving high compression ratios while maintaining a low compression cost. The Boost Graph Library and the Parallel Boost Graph Library both provide appealing features for working with graphs, with the latter library notably having interoperability with MPI, Graphviz, and METIS [22], [23].

2.3 Improvements Specific to Attack Graph Generation

As a means of improving scalability of attack graphs, the authors of [3] present a new representation scheme. Traditional attack graphs encode the entire network at each state, but the representation presented by the authors uses logical statements to represent a portion of the network at each node. This is called a logical attack graph. This approach led to the reduction of the generation process to quadratic time and reduced the number of nodes in the resulting graph to $\mathcal{O}(n^2)$. However, this approach does require more analysis for identifying attack vectors. Another approach presented by the authors of [24] represents a description

of systems and their qualities and topologies as a state, with a queue of unexplored states. This work was continued by the authors of [25] by implementing a hash table among other features. Each of these works demonstrates an improvement in scalability through refining the desirable information output.

Another approach for generation improvement is through parallelization. The authors of [25] leverage OpenMP to parallelize the exploration of a FIFO queue. This parallelization also includes the utilization of OpenMP’s dynamic scheduling. In this approach, each thread receives a state to explore, where a critical section is employed to handle the atomic functions of merging new state information while avoiding collisions, race conditions, or stale data usage. The authors measured a 10x speedup over the serial algorithm. The authors of [26] present a parallel generation approach using CUDA, where speedup is obtained through a large number of CUDA cores. For a distributed approach, the authors of [27] present a technique for utilizing reachability hyper-graph partitioning and a virtual shared memory abstraction to prevent duplicate work by multiple nodes. This work had promising results in terms of limiting the state-space explosion and speedup as the number of network hosts increases.

CHAPTER 3

UTILITY EXTENSIONS TO THE RAGE ATTACK GRAPH GENERATOR

3.1 Path Walking

Due to the large-scale nature of attack graphs, analysis can become difficult and time-consuming. With some graphs reaching millions of states and edges, analyzing the entire graph can be overwhelmingly complex. As a means of simplifying analysis, a potential strategy could be to consider only small subsets of the graph at a time, rather than feeding the entire graph into an analysis algorithm. This allows users to focus only on specific states or items within the graph without needing to analyze the entirety of the graph. To aid in this effort, a path walking feature was implemented as a separate program, and has two primary modes of usage. The goal of this feature is to output a subset of the graph that includes all possible paths from the root state to a designated state. The first mode is a manual mode, where a user can input the desired state to walk to, and the program will output a separate graph of all possible paths to the specified state. The second mode is an automatic mode, where the program will output separate subgraphs to all states in the network that have qualities of “*compliance_vio = true*”, “*compliance_vios > 0*”, or any other quality that can be specified by the user. The automatic mode can produce multiple subgraphs simultaneously if multiple states contain the quality being examined, and these subgraphs can then be separately input into an analysis program.

Figure 3.1 demonstrates an output of the Path Walking feature when walking to state 14. In this figure, the primary observable feature is that the graph was reduced from 16 states to 6 states, and from 32 edges to 12 edges. The amount of reduction from the original graph to the subset varies on the overall connectivity of the original attack graph.

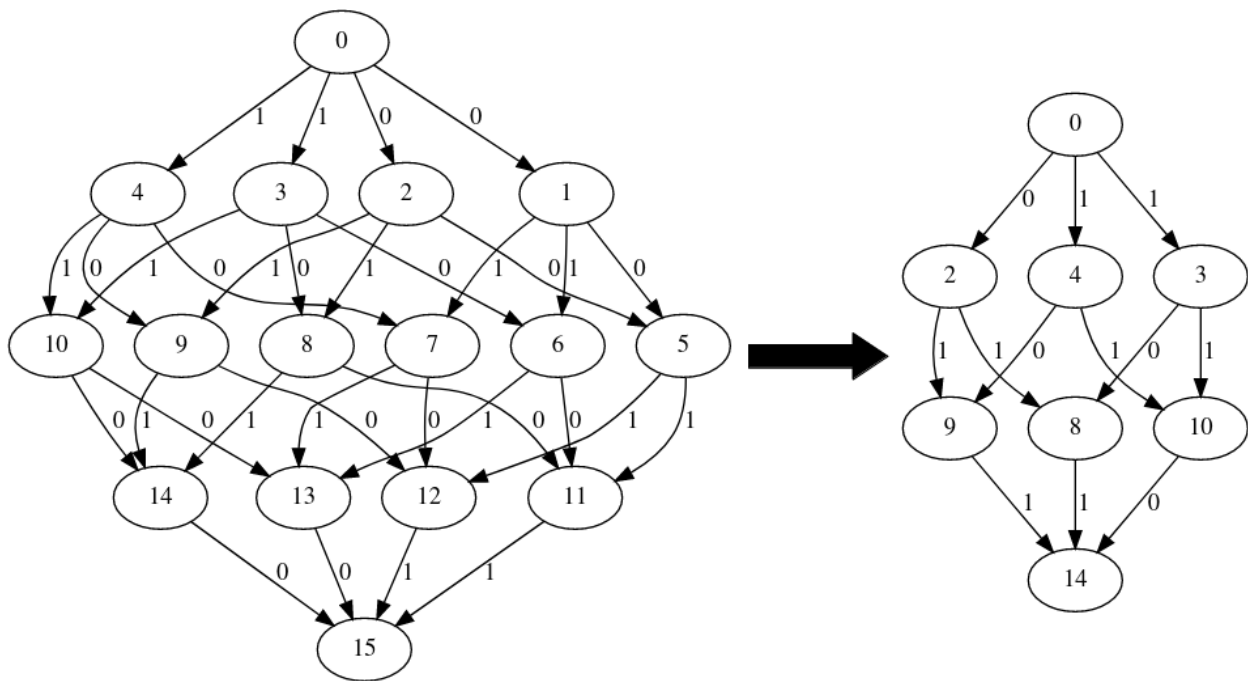


Figure 3.1: Path Walking to State 14

3.2 Color Coding

As a visual aid for analysis purposes, color coding was another feature implemented as a postprocessing tool for RAGE. When viewing the output graph of RAGE, all states are visibly identical in appearance apart from number of edges, edge IDs, and state IDs. To allow for visual differentiation, color coding can be enabled in the run script. Color coding currently functions by working through the graph output text file, but it can be extended to read directly from the PostgreSQL database instead. The feature scans through the output file, and locates states that have “*compliance_vios* > 0” or “*compliance_vio* = true”. For states that meet these properties, the color coding feature will add a color to the Graphviz DOT [28] file through the [*color* = *COL*] attribute for the given node, where *COL* is assigned based on severity. For this version of color coding, severity is determined by the total number of compliance violations a node has, but future versions can alter the severity measure through alternative means. Figure 3.2 displays an example graph that leverages color coding to easily identify problem states. Each shaded state represents a problem state, where the system

contains items that are out of compliance or in violation of a regulation. For this specific example, each shaded state represents a state where the vehicle is in need of maintenance on a component, as it has exceeded the recommended time since the last maintenance.

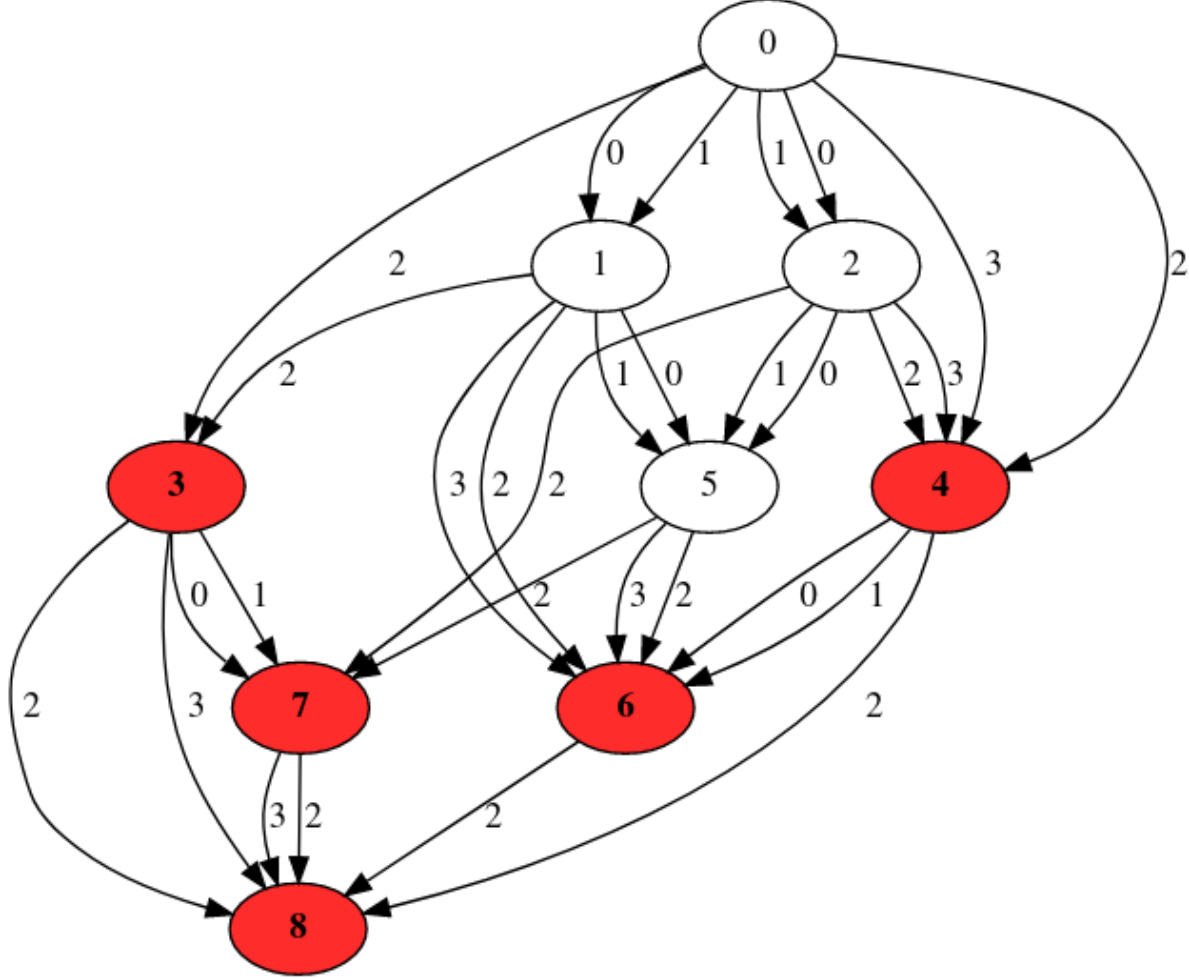


Figure 3.2: Color Coding a Small Network Based on Violations

3.3 Compound Operators

Many of the graphs previously generated by RAGE comprise of states with features that can be fully enumerated. In many of these generated graphs, there was an established set of qualities that was used, with an established set of values. These typically have included “*compliance_vio = true/false*”, “*root = true/false*”, or other general “*true/false*” values

or “*version = X*” qualities. To expand on the types and complexities of graphs that can be generated, compound operators have been added to RAGE. When updating a state, rather than setting a quality to a specific value, the previous value can now be modified by an amount specified through standard compound operators such as $+=$, $-=$, $*=$, or $/=$. Previous work on an attack graph generator included the implementation of compound operators, as seen by the author of [29]. However, this work was conducted on the previous iteration of an attack graph generator written in Python. This attack graph generator has since been rewritten in C++ by the author of [13], and compound operators have not been included in the latest version of RAGE.

The work conducted by the author of [13] when designing the software architecture of RAGE included specifications for a quality encoding scheme. As the author discusses, qualities have four fields, which include the asset ID, attributes, operator, and value. The operator field is 4 bits, which allows for a total of 16 operators. Since the only operator in use at the time was the “ $=$ ” operator, the addition of four compound operators does not surpass the 16 operator limit, and no encoding scheme changes were necessary. This also allows for additional compound operators to be incorporated in the future.

A few changes were necessary to allow for the addition of compound operators. Before the generation of an attack graph begins, all values are stored in a hash table. For previous networks generated by RAGE, this was not a concern since all values could be fully enumerated and all possible values were known. When using compound operators however, not all values can be fully known. The task of approximating which exploits will be applicable and what absolute minimum or maximum value bounds will be prior to generation is difficult, and not all values can be enumerated and stored into the hash table. As a result, real-time updates to the hash table needed to be added to the generator. The original key-value scheme for hash tables relied on utilizing the size of the hash table for values. Since the order in which updates happen may not always remain consistent (and is especially true in distributed computing environments), it is possible for states to receive different hash values with the original hashing scheme. To prevent this, the hashing scheme was adjusted so that

the new value of the compound operator is inserted into the hash table values if it was not found, rather than the size of the hash table. Previously, there was no safety check for the hash table, so if the value was not found, the program would end execution. The assertion that the new value can be inserted into the hash table is safe to make, since compound operators are conducted on numeric values, and matches the numeric type of the hash table.

Other changes involved updating classes (namely the Quality, EncodedQuality, ParameterizedQuality, NetworkState, and Keyvalue classes) to include a new member for the operator in question. Auxiliary functions related to this new member, such as prints and getters, were also added. In addition, preconditions were altered to include operator overloads to check the asset identifier, quality name, and quality values for the update process.

3.4 Relational Operators

As discussed in Section 3.3, many of the graphs previously generated by RAGE comprise of states with an established set of qualities and values. These typically have included “*compliance_vio = true/false*”, “*root = true/false*”, or other general “*true/false*” values or “*version = X*” qualities. To further expand the dynamism of graph generation, it is important to distinguish when a quality has a value that satisfies a relational comparison to an exploit. An example application for attack graphs can be seen through CVE-2019-10747, where “set-value is vulnerable to Prototype Pollution in versions lower than 3.0.1” [30]. An example compliance graph application using the aforementioned car example can be seen in the Toyota Corolla Maintenance Schedule, which states an engine coolant replacement should be conducted after 24,000 miles. Prior to the implementation of relational operators, to determine whether this exploit was applicable to a network state, multiple exploit qualities must be enumerated for all versions prior to 3.0.1. This would mean that the exploit needed to check if *version=3.0.0*, or *version=2.0.0*, or *version=1.0.0*, or *version=0.4.3*, etc. For the compliance graph exploit check, this could lead to even worse scaling where checks needed to be conducted at a much more granular level like *engine_coolant_miles=24001*, or *engine_coolant_miles=24002*, or *engine_coolant_miles=24003*, etc. This becomes increasingly

tedious when there are many checks to perform, and this not only reduces readability, but is also more prone to human error when creating the exploit files. Relational operators work to alleviate these difficulties.

Similar to the compound operators discussed in Section 3.3, relational operators were also implemented in a previous attack graph generator, and were implemented as a part of the works performed by [29]. This implementation was in the previous Python generator tool, and relational operators were not a feature of RAGE. The implementation of relational operators in RAGE allows users to leverage the performance benefits of the C++ generator and have the ability for parallelization, while providing ease-of-use for exploit file creation.

To implement the relational operators, operator overloads were placed into the `Quality` class. At the time of writing, the following are implemented: `==`, `<`, `>`, `≤`, `≥`. However, these operators do not take up room in the encoding scheme, so additional operators can be freely implemented as needed. The overloads ensure that the `Quality` asset IDs and `Quality` names match, and then compares the `Quality` values based on the operator in question.

3.5 Intermediate Database Storage

3.5.1 Introduction to Intermediate Database Storage

Chapter 2 and the author of [13] discuss the challenges of attack graph generation in regards to its scalability. Specifically, the author of [13] displays results from generations based on 11 assets and 11 exploits that lead to 14,200 total states. Generating an attack or compliance graph based on a large network with a multitude of assets and involving a more thorough exploit or compliance violation checking will prevent the entire graph from being stored in memory as originally designed. This Section discusses the challenges of graph generation in regards to memory, and a solution through the implementation of intermediate database storage.

3.5.2 Memory Constraint Difficulties

Previous works with RAGE have been designed around maximizing performance to

limit the longer runtime caused by the state space explosion, such as the works seen by the authors of [13], [25], and [31]. To this end, the output graph is contained in memory during the generation process to minimize disk writing and reading. This also allows for leveraging the performance benefits of memory operations, since graph computation relies less on processor speed and more on data dependency complexity, parallelism coarseness, and memory access time [5], [15], [14]. The author of [13] does incorporate PostgreSQL as an initial and final storage mechanism to write the starting and resulting graph information, but no intermediate storage is otherwise conducted.

While the design decision to not use intermediate storage maximizes performance for graph generation, it introduces a few complications. When generating large graphs, the system runs the risk of running out of memory. This typically does not occur when generation is conducted on small graphs, and is especially true when relatively small graphs are generated on an HPC system with substantial amounts of memory. However, when running on local systems or when the graph is large, memory can quickly be depleted due to state space explosion. The memory depletion is due to two primary memory consumption points: the frontier which contains all of the states that still need to be explored, and the graph instance which holds all of the states and their information, as well as all of the edges.

The frontier quickly becomes a problem point with large graphs that have a large height value, and contain many layers before reaching leaf nodes. During the generation process, RAGE works on a Breadth-First Search approach, and new states are continuously discovered each time a state from the frontier is explored. In almost all cases, this means that for every state that is removed from the frontier, several more are added, leading to an ever-growing frontier that can not be adequately reduced for large networks. Simultaneously, the graph instance continues to grow as states are explored. When the network contains numerous assets, each with their own large sets of qualities, the size of each state becomes noticeably larger. With some graphs containing millions of nodes and billions of edges like those mentioned by the authors of [5], it becomes increasingly unlikely that the graph can be fully contained within system memory.

3.5.3 *Maximizing Performance with Intermediate Database Storage*

Rather than a static implementation of storing to the database on disk at a set interval or a set size, the goal was to dynamically store to the database only when necessary. This would allow for proper utilization of systems with greater memory, and would reduce fine-tuning of a maximum size variable before database writes on different systems. Since there is an associated cost with preparing the writes to disk, the communication cost across nodes, the writing to disk itself, and a cost for retrieving items from disk, it is desirable to store as much in memory for as long as possible and only write when necessary. When running RAGE, a new argument can be passed (*-a <double>*) to specify the amount of memory the tool should use before writing to disk. This argument is a value between 0 and 0.99 to specify a percentage. This double is immediately reduced by 10%. For instance, if 0.6 is passed, it is immediately reduced to 0.5. This acts as a buffer for PostgreSQL. Since queries will consume a variable amount of memory through parsing or preparation, an additional 10% is saved as a precaution. This can be changed later as needed or desired for future optimizations. Specific to the graph data, the statement is made that the frontier is allowed to consume half of the allocated memory, and that the graph instance is allowed to consume the other half.

To decide when to store to the database instead of memory, two separate checks are made. The first check is for the frontier. If the size of the frontier consumes equal to or more than the allowed allocated memory, then all new states are stored into a new table in the database called “unexplored states”. Each new state from this point forward is stored in the table, regardless of if room is freed in the frontier. This is to ensure proper ordering of the FIFO queue. The only time new states are stored directly into the frontier is when the unexplored states table is empty. Once the frontier has been completely emptied, new states are then pulled from the database into the frontier. To pull from the database, the parent loop for the generator process has been altered. Instead of a while loop for when the frontier is not empty, it has been adjusted to when the frontier is not empty or the unexplored states table is not empty. Due to C++ using short-circuit evaluation where

the first argument is completely evaluated before processing the second, some performance is gained. The performance gained is due to not having to pass a SQL statement to disk to check the size of the unexplored states table unless the frontier is empty. The original generation design stored new states into the frontier during the critical section to avoid testing on already-explored states. To follow this design decision, writing new states to the database is also performed during the critical section.

For the graph instance, a check in the critical section determines if the size of the graph instance consumes more than its allocated share of the memory. If it does, the edges, network states, and network state items are written to the database, and are then removed from memory.

However, a new issue arose with database storage. The original design was to save staging, preparation, and communication cost by writing all the data in one query (as in, writing all of the network states in one query, all the network state items in one query, and all the edges in one query). While this was best option in terms of performance, it was also not feasible when the amount of data to store was large in relation to system memory. Building the SQL queries themselves quickly began depleting the already constrained memory with large storage requests. As a result, the storage process would consume too much memory and crash the generator tool. To combat this, all queries had to be broken up into multiple queries. As previously mentioned, an extra 10% buffer was saved for the storage process. SQL query strings are now built until they consume the 10% buffer, where they are then processed by PostgreSQL, cleared, and the query building process resumes.

3.5.4 *Portability*

The intermediate database storage is greatly advantageous in increasing the portability of RAGE across various systems, while still allowing for performance benefits. By allowing for a user-defined argument, users can safely assign a value that allows for other processes and for the host OS to continue their workloads. While the “total memory” component currently utilizes the Linux *sysconf()* function, this is not rigid and is easily adjustable.

When working on an HPC cluster, using this function could lead to difficulties since multiple users may be working on the same nodes, which prevents RAGE from fully using all system memory. This could be prevented by using a job scheduler argument such as Slurm’s “--exclusive” option, but this may not be desirable. Instead, a user could pass in the amount of total memory to use (and can be reused from a job scheduler’s memory allocation request option), and the intermediate database storage process would function in the same fashion.

CHAPTER 4

SYNCHRONOUS FIRING

4.1 Introduction

One main appeal of attack graphs and compliance graphs are their exhaustiveness. The ability to generate all permutations of attack chains or to generate all possible ways a system can fall out of compliance is a valuable feature. The disadvantage of this approach is that the generation of the final graph increases in time, as does the analysis. One other disadvantage is that this exhaustiveness can produce states that are not actually attainable. When a system has assets that have inseparable features, the generation process forcibly separates features to examine all permutations, since the generation process only modifies one quality at a time. One example of an inseparable feature is time. If two different assets are identical and no constraints dictate otherwise, the two assets cannot proceed through times at different rates.

For example, if two cars were manufactured at the same moment, one of these cars cannot proceed multiple time steps into the future while the other remains at its current time step; each car must step through time at the same rate. However, the generation of attack graphs and compliance graphs examines the possibilities that one car ages by one time step, while the other car does not, or vice versa. This results in an attack graph that can be seen in Figure 4.1, which is a partial attack graph showing the separation of the time feature. All shaded states are considered unattainable, since all of these states comprise of assets that have advanced time at different rates. It is noticeable that not only are the unattainable states themselves a wasteful generation, but they also lead to the generation of even more unattainable states that will then also be explored. A better procedure for a

generation process similar to this example is to have a single state transition that updates assets with an inseparable feature simultaneously.

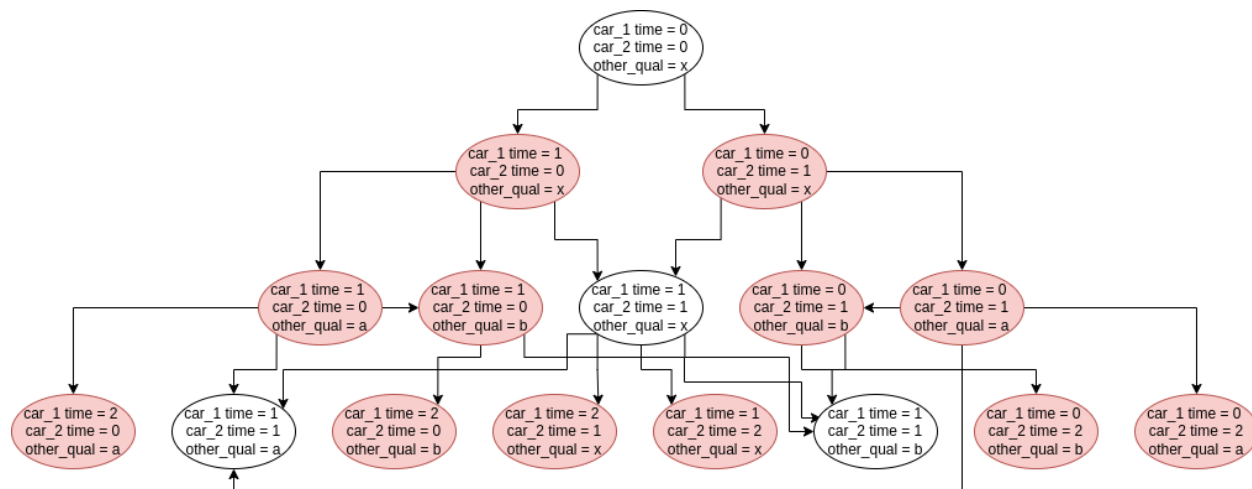


Figure 4.1: A network without Synchronous Firing generating infeasible states

Post-processing is one option at removing the unattainable states. This process would simplify and reduce the time taken for the analysis process, but the generation process would still suffer from generating and exploring the unattainable states. Instead, a new feature called synchronous firing can be used to prevent the generation of these states. The goal of the synchronous firing feature is to prevent the generation of unattainable states, while also not incurring a greater computational cost. This Chapter will discuss the development of this feature, discuss its references and inspiration in literature, and examine the results when using this feature in applicable networks.

4.1.1 Related Synchronous Firing Work

A form of synchronous firing is discussed by the author of [32], where it is described as grouped exploits. The functionality discussed by the author is similar: firing an exploit should be performed on all possible assets simultaneously. This was also described as synchronizing multiple exploits. The methodology is similar to the one implemented in this work, but there are notable differences. The first, is that the work performed by the author

of [32] utilizes global features with group features. Using the simultaneous exploit firing necessitated a separation of global and group features, and grouped exploits could not be performed on exploits that could be applicable to both sets. A second difference is that there is no consistency checking in the work by the author of [32], which could lead to indeterminate behavior or race conditions unless additional effort was put into encoding exploits to use precondition guards. A third difference is that the work of [32] could still lead to a separation of features. The grouped exploit feature would attempt to fire all exploits on all applicable assets simultaneously, but if some assets were not ready or capable to fire, these assets would not proceed with the exploit firing but the applicable assets would. The last difference is that the work by the author of [32] was developed in Python, since that was the language of the generator of the tool at the time. The work by the author of [13] led to new development of RAGE in C++ for performance enhancements, so the synchronous firing feature in this new work was likewise developed in C++.

4.2 Necessary Alterations and Additions

For the implementation of the synchronous firing feature, there were four primary changes and additions that were necessary. The first is a change in the lexical analyzer, the second involves multiple changes to PostgreSQL, the third is the implementation of compound operators (as discussed in Section 3.3), and lastly is a change in the graph generation process. This Section will consist of subsections discussing the development of these four alterations.

4.2.1 *GNU Bison and Flex*

The work conducted by the author of [13] included the introduction of GNU Bison and GNU Flex into RAGE. The introduction of Bison and Flex allows for an easily modifiable grammar to adjust features, the ability to easily update parsers since Bison and Flex are built into the build system, and increases portability since Flex and Bison generate standard C. For the development of the synchronous firing feature, a similar approach was taken to

that of the work performed by the author of [32] with the exploit keywords. However, rather than having both global and group keywords, this work only incorporates the group keyword to prevent a few of the difficulties discussed in Section 4.1.1. The new “group” keyword is intended to be used when creating the exploit files. The design of exploits in the exploit file is developed as:

```
<exploit> ::= <group name> "group" "exploit" <identifier> , (<parameter-list>)=
```

where the “<group name>” identifier and “group” keyword is optional. An example of an exploit not utilizing the group feature is:

```
exploit brake_pads(2015_Toyota_Corolla_LE)=
```

and an example of an exploit utilizing the group feature is:

```
time group exploit advance_month(all_applicable)=
```

To implement the keyword recognition and group name parsing, a few changes were made, where the intention was to detect the usage of the “group” keyword, and have the lexical analyzer code return to the parser implementation file to alert of the presence of the “GROUP” token. The new token is of type string with the name “GROUP”, and it is comprised of a leading “IDENTIFIER” of type string or integer token, followed by the “GROUP” token. This new token also required changes to the processing of the “exploit” keyword. If the group keyword is not detected, the exploit has a group of name “null”. If the group keyword is detected, then the leading IDENTIFIER is parsed, and the exploit is assigned to a group with the parsed name. Various auxiliary functions were also adjusted to include (for instance) support for printing the groups of each exploit. Figure 4.2 illustrates the incorporation of this feature into Bison, Flex, and the overall program.

4.2.2 PostgreSQL

As seen in Figure 4.2, Bison and Flex feed into the Model Database. With the addition of a new group identifier and the group keyword, minor alterations were needed to

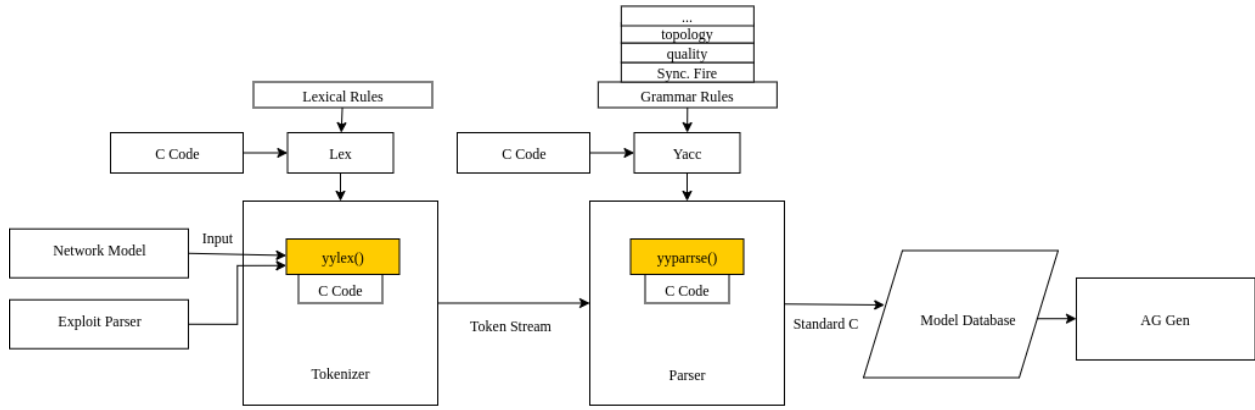


Figure 4.2: Inclusion of Synchronous Firing into GNU Bison, GNU Flex, and the overall program

ensure compatibility with the PostgreSQL database. One adjustment was to alter the exploit table in the SQL schema to include new columns of type “TEXT”. The second adjustment was to update the SQL builder functions. This included updating the related functions such as exploit creations, exploit parsing, database fetching, and SQL string builders to add additional room for the group identifier.

4.2.3 Compound Operators

While not strictly necessary, compound operators greatly simplify the exploit file creation process. For example, implementing time as a feature into the tool without compound operators would increase its difficulty substantially. For each time interval, a separate exploit would need to be created, with time flags to indicate the current time. If time was increased monthly for a year, 12 different exploits would need to be created, with flags to ensure that time jumps did not occur. Updating qualities without compound operators also relied on flags, and clever, but convoluted methods for incrementing values. Instead, compound operators were implemented, and this addition was discussed in Section 3.3.

4.2.4 Graph Generation

The implementation of synchronous firing in the graph generation process relies on a map to hold the fired status of groups. Previously, each iteration of the applicable exploit

vector loop generated a new state. With synchronous firing, all assets should be updating the same state, rather than each independently creating a new state. To implement this, each iteration of the applicable exploit vector checks if the current loop element is in a group and if that group has fired. If the element is in a group, the group has not been fired, and all group members are ready to fire, then all group members will loop through an update process to alter the single converged state. Otherwise, the loop will either continue to the next iteration if group conditions are not met, or will create a single state if it is not in a group. Figure 4.3 displays the synchronous fire approach.

4.3 Experimental Networks and Results

All data was collected on a 13 node cluster, with 12 nodes serving as dedicated compute nodes, and 1 node serving as the login node. Each compute node has a configuration as follows:

- OS: CentOS release 6.9
- CPU: Two Intel Xeon E5-2620 v3
- Two Intel Xeon Phi Co-Processors
- One FPGA (Nallatech PCIE-385n A7 Altera Stratix V)
- Memory: 64318MiB

All nodes are connected with a 10Gbps Infiniband interconnect.

4.3.1 *Experimental Networks*

The example networks for testing the effectiveness of synchronous firing follow the compliance graph generation approach. These networks analyze two assets, both of which are identical 2006 Toyota Corolla cars with identical qualities. The generation examines both cars at their current states, and proceeds to advance in time by a pre-determined amount, up to a pre-determined limit. Each time increment updates each car by an identical amount of

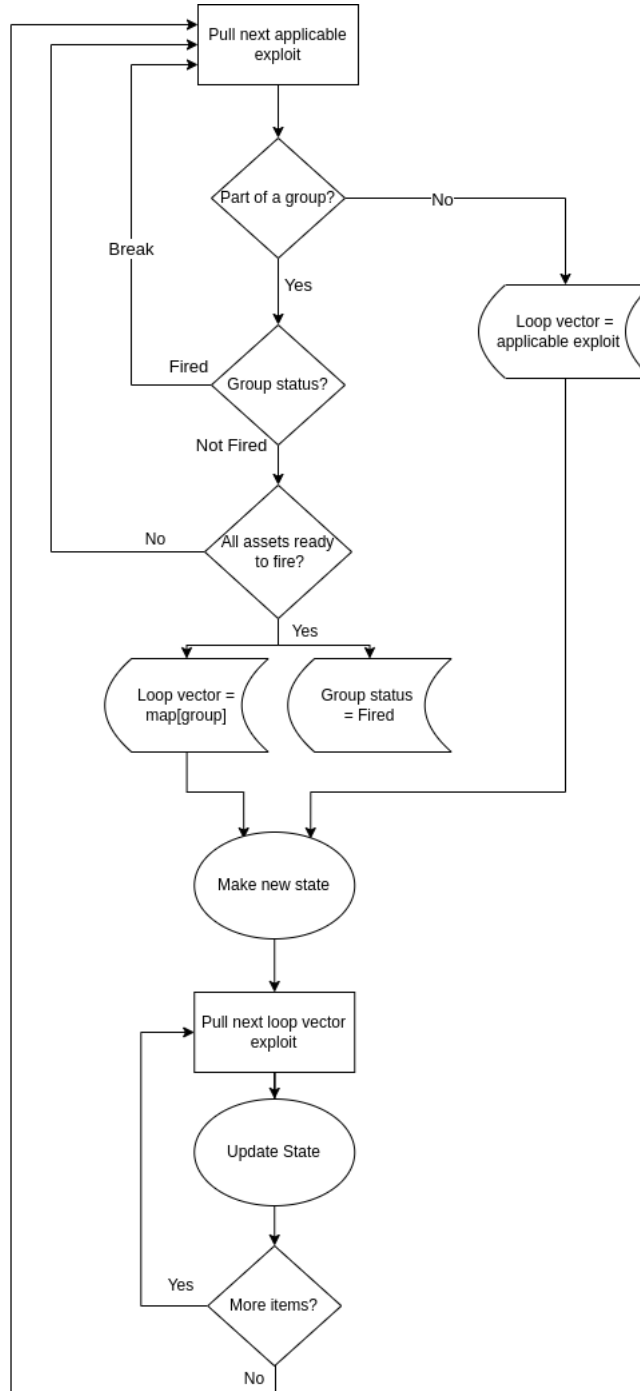


Figure 4.3: Synchronous Firing in the Graph Generation Process

mileage. During the generation process, it is determined if a car is out of compliance either through mileage or time since its last maintenance in accordance with the Toyota Corolla Maintenance Schedule manual.

In addition, the tests employ the use of “services”, where if a car is out of compliance, it will go through a correction process and reset the mileage and time since last service. Each test varies in the number of services used. The 1 Service test only employs one service, and it is dedicated to brake pads. The 2 Service test employs two services, where the first service is dedicated to the brake pads, and the second is for exhaust pipes. This process extends to the 3 and 4 Service tests. The testing information is as follows:

- All tests ran for 12 months, with time steps of 1 month.
- All tests had the same number of compliance checks: brake pads, exhaust pipes, vacuum pumps, and AC filters.
- There were 10 base exploits, and an additional 4 exploits were individually added in the form of services for each test.
- All tests used the same network model.
- All tests used the same exploit file, with the exception of the “group” keyword being present in the synchronous firing testing.
- Services must be performed prior to advancing time, if services are applicable.
- Graph visualization was not timed. Only the generation and database operation time was measured.

4.3.2 Results

Using the testing setup described in Section 4.3.1 on the platform described at the beginning of Section 4.3, results were collected in regards to the effect of synchronous firing on both state space and runtime. These results can be seen in Figures 4.4 and 4.5. Both figures show a decrease in the number of states and a decrease in the runtime when synchronous firing

is utilized. Since synchronous firing prevents the generation of unattainable states, there is no meaningful information loss that occurs in the graphs generated with the synchronous firing feature. It is also noteworthy that even for the smallest graph generated (which only resulted in 394 states with synchronous firing enabled), the computation cost of additional work for group vectors was justified by the decrease in runtime. Since the resulting number of states was also reduced, there will be increased justification for the synchronous firing approach due to a reduced runtime for the analysis process. Further elaboration is seen in Section 6.2, but it is hypothesized that the synchronous firing approach will lead to an increased runtime reduction and state space reduction when more assets exist in the network due to the increased number of unattainable state permutations.

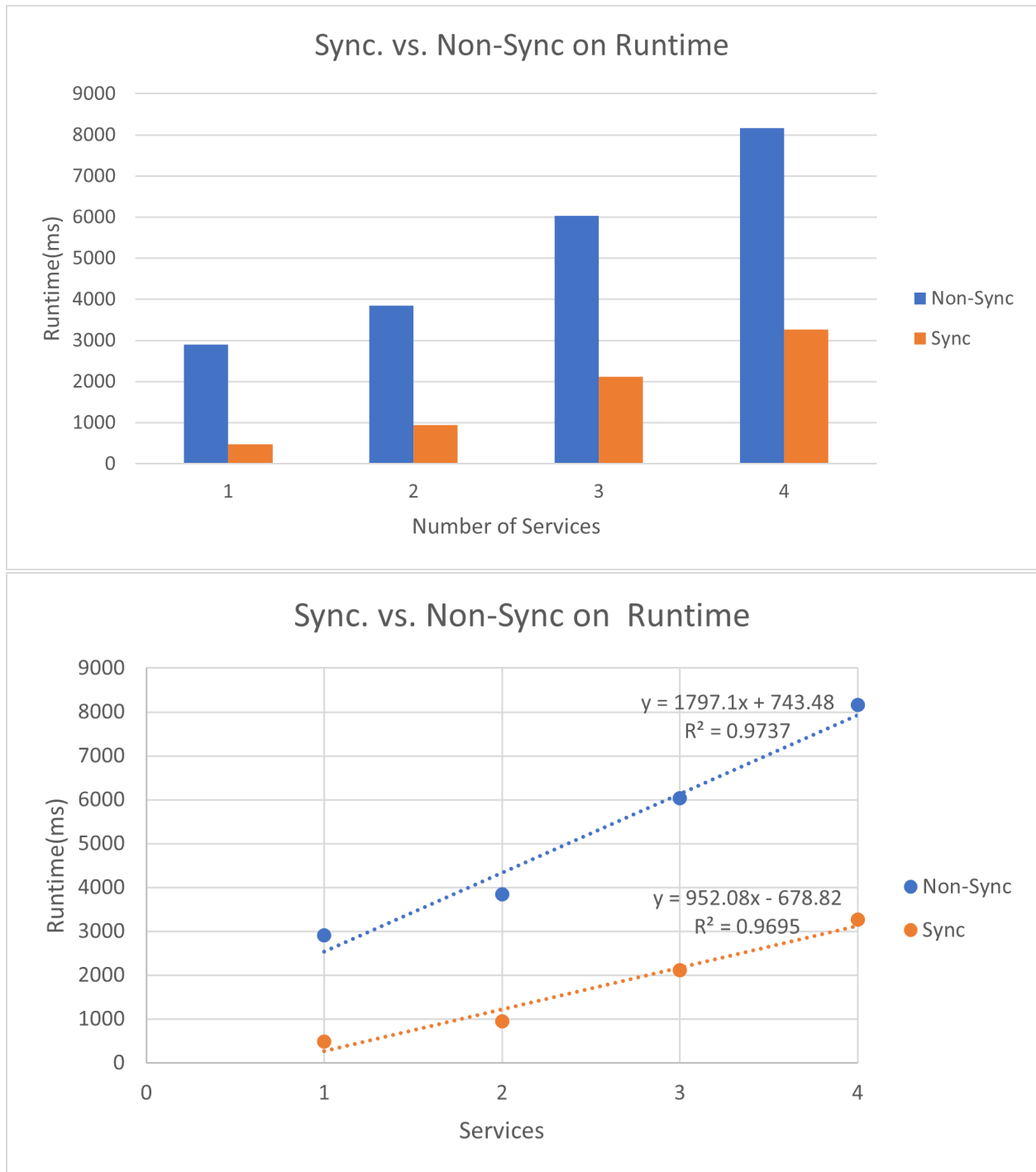


Figure 4.4: Bar Graph and Line Graph Representations of Synchronous Firing on Runtime

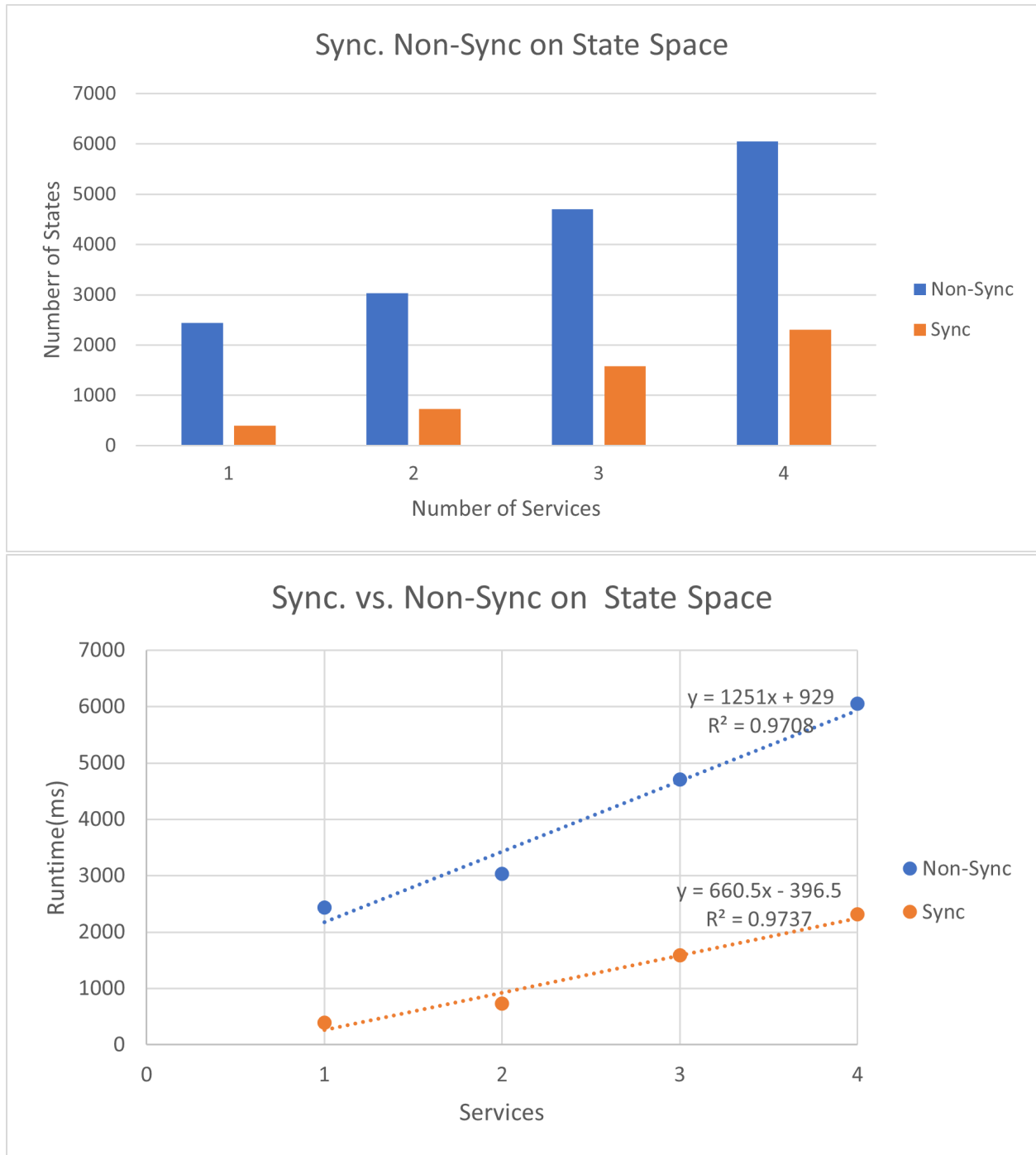


Figure 4.5: Bar Graph and Line Graph Representations of Synchronous Firing on State Space

CHAPTER 5

PARALLELIZATION USING MESSAGE PASSING INTERFACE

5.1 Introduction to MPI Utilization for Attack and Compliance Graph Generation

Previous works for graph generation, and specifically for attack graph generation, have seen promising results as discussed in Sections 2.2 and 2.3. This work attempts to further those efforts and extend RAGE to function on distributed computing environments to take advantage of the increased computing power using message-passing. As mentioned by the author of [33], MPI is the most widely used message-passing API, and this work intended to utilize an API that was not only familiar and accessible, but versatile and powerful for parallelizing RAGE for distributed computing platforms. This Chapter discusses two approaches for parallelism: task parallelism in Section 5.3, and data parallelism in Section 5.4. All approaches in this work utilize OpenMPI for the MPI implementation.

5.2 Necessary Components

5.2.1 *Serialization*

In order to distribute workloads across nodes in a distributed system, various types of data will need to be sent and received. Support and mechanisms vary based on the MPI implementation, but most fundamental data types such as integers, doubles, characters, and Booleans are incorporated into the MPI implementation. While this does simplify some of the messages that need to be sent and received in the MPI approaches of attack and compliance graph generation, it does not cover the vast majority of them when using RAGE.

RAGE implements many custom classes and structs that are used throughout the generation process. Qualities, topologies, network states, and exploits are a few such examples. Rather than breaking each of these down into fundamental types manually, serialization functions are leveraged to handle most of this. RAGE already incorporates Boost graph libraries for auxiliary support, so this work extended this further to utilize the serialization libraries also provided by Boost. These libraries also include support for serializing all STL classes, and many of the RAGE classes have members that make use of the STL classes. One additional advantage of the Boost library approach is that many of the RAGE classes are nested. For example, the `NetworkState` class has a member vector of `Quality` classes, and the `Quality` class has a `Keyvalue` class as a member. When serializing the `NetworkState` class, Boost will recursively serialize all members, including the custom class members, assuming they also have serialization functions.

When using the serialization libraries, this work opted to use the intrusive route, where the class instances are altered directly. This was preferable to the non-intrusive approach, since the class instances were able to be altered with relative ease, and many of the class instances did not expose enough information for the non-intrusive approach to be viable.

5.3 Tasking Approach

5.3.1 Introduction to the Tasking Approach

The high-level overview of the attack and compliance graph generation process can be broken down into six main tasks. These tasks are described in Figure 5.1. Prior works such as that seen by the authors of [25], [26], and [27] work to parallelize the graph generation using OpenMP, CUDA, and hyper-graph partitioning. This approach, however, utilizes Message Passing Interface (MPI) to distribute the six identified tasks of RAGE to examine the effect on speedup, efficiency, and scalability for attack and compliance graph generation.

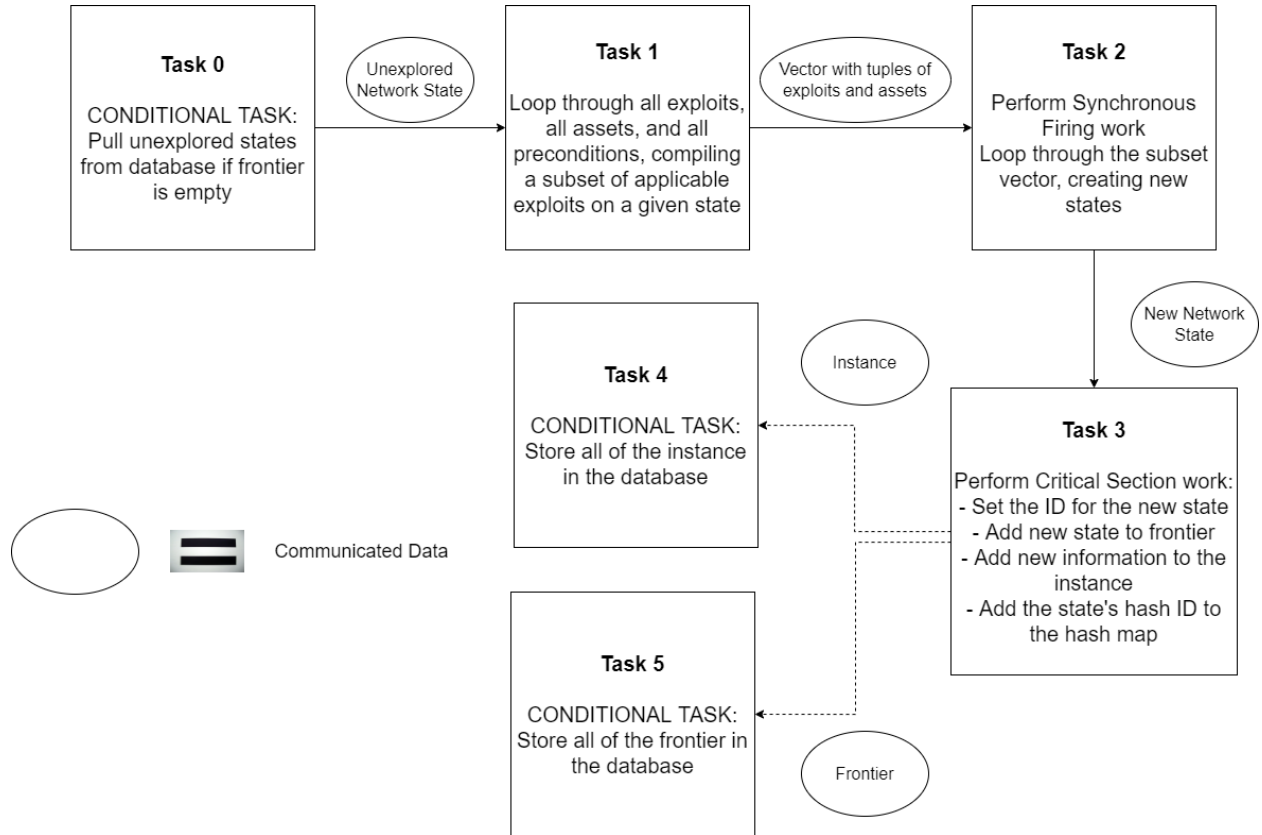


Figure 5.1: Task Overview of the Attack and Compliance Graph Generation Process

5.3.2 Algorithm Design

The design of the tasking approach is to leverage a pipeline structure with the six tasks and MPI nodes. After its completion, each stage of the pipeline will pass the necessary data to the next stage through various MPI messages, where the next stage's nodes will receive the data and execute their tasks. The pipeline is considered fully saturated when each task has a dedicated node solely for executing work for that task. When there are less nodes than tasks, some nodes will process multiple tasks. When there are more nodes than tasks, additional nodes will be assigned to Tasks 1 and 2. Timings were collected in the serial approach for various networks that displayed more time requirements for Tasks 1 and 2, with larger network sizes requiring vastly more time to be taken in Tasks 1 and 2. As a result, additional nodes are assigned to Tasks 1 and 2. Node allocation can be seen in Figure 5.2. In this Figure, “world.size()” is an integer value representing the number of nodes used in the program, and “num_tasks” is an integer value representing the number of tasks used in the pipeline. By using a variable for the number of tasks, it allows for modular usage of the pipeline, where tasks can be added and removed without needing to change any allocation logic work; only communication between tasks may need to be modified, and the allocation can be adjusted relatively simply to include new tasks.

For determining which tasks should be handled by the root node, a few considerations were made, where minimizing communication cost and avoiding unnecessary complexity were the main two considerations. In the serial approach, the frontier queue was the primary data structure for the majority of the generation process. Rather than using a distributed queue or passing multiple sub-queues between nodes, the minimum cost option is to pass states individually. This approach also assists in reducing the complexity. Managing multiple frontier queues would require duplication checks, multiple nodes requesting data from and storing data into the database, and devising a strategy to maintain proper queue ordering, all of which would also increase the communication cost. As a result, the root node will be dedicated to Tasks 0 and 3.

Task	Node Rank(s) Allocated
0	0
1	$[1, n_1]$
2	$[(n_1 + 1), n_2]$
3	0
4	n_3
5	n_4

$$n_1 = \begin{cases} 1, & \text{world.size() } \leq \text{num_tasks} \\ 1 + \lceil \frac{\text{world.size()} - \text{num_tasks}}{2} \rceil, & \text{otherwise} \end{cases}$$

$$n_2 = \begin{cases} 2n_1, & \text{world.size() \% 2 = 0} \\ 2n_1 - 1, & \text{otherwise} \end{cases}$$

$$n_3 = n_2 + 1$$

$$n_4 = n_3 + 1$$

Figure 5.2: Node Allocation for each Task

5.3.2.1 Communication Structure: The underlying communication structure for the tasking approach relies on a pseudo-ring structure. As seen in Figure 5.2, nodes n_2 , n_3 , and n_4 are derived from the previous task’s greatest node rank. To keep the development abstract, a custom send function checks the world size (“world.size()”) before sending. If the rank of the node that would receive a message is greater than the world size and therefore does not exist, the rank would then be “looped around” and corrected to fit within the world size constraints. After the rank correction, the MPI Send function was then invoked with the proper node rank.

5.3.2.2 Task 0: Task 0 is performed by the root node, and is a conditional task; it is not guaranteed to be executed at each pipeline iteration. Task 0 is only executed when the frontier is empty, but the database still holds unexplored states. This occurs when there are memory constraints, and database storage is performed during execution to offload the demand, as discussed in Section 3.5. After the completion of Task 0, the frontier has a state popped, and the root node sends the state to n_1 . If the frontier is empty, the root node sends the finalize signal to all nodes.

5.3.2.3 Task 1: Task 1 begins by distributing the workload between nodes based on the local task communicator rank. Rather than splitting the exploit list at the root node and sending sub-lists to each node allocated to Task 1, each node checks its local communicator rank and performs a modulo operation with the number of nodes allocated to determine whether it should proceed with the current iteration of the exploit loop. Since the exploit list is static, each node has the exploit list initialized prior to the generation process, and communication cost can be avoided from sending sub-lists to each node. Each node in Task 1 works to compile a reduced exploit list that is applicable to the current network state. A breakdown of the Task 1 distribution can be seen in Figure 5.3.

Once the computation work of Task 1 is completed, each node must send their compiled applicable exploit list to Task 2. Rather than merging all lists and splitting them back out in Task 2, each node in Task 1 will send an applicable exploit list to at most one node

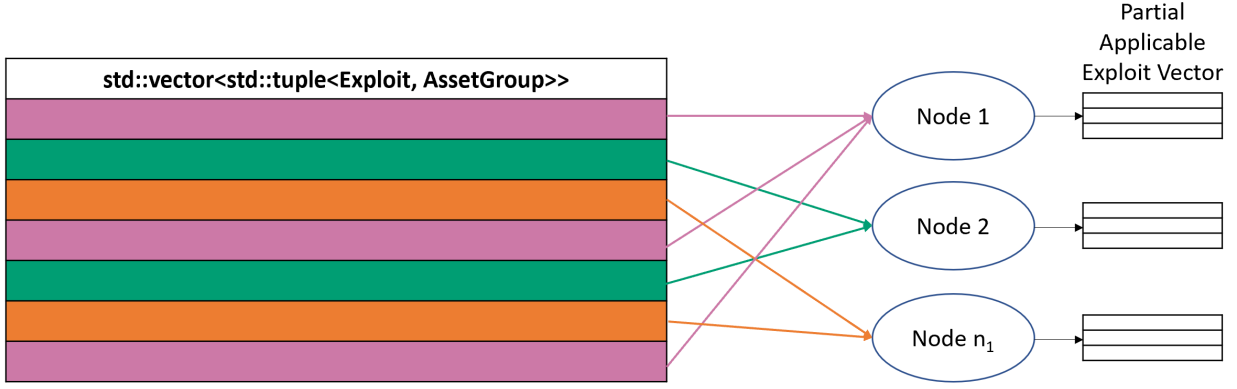


Figure 5.3: Data Distribution of Task One

allocated to Task 2. Based on the allocation of nodes seen in Figure 5.2, there are 2 potential cases: the number of nodes allocated to Task 1 is equal to the number of nodes allocated to Task 2, or the number of nodes allocated to Task 1 is one greater than the number of nodes allocated to Task 2. For the first case, each node in Task 1 sends the applicable exploit list to its global rank+ n_1 . This case can be seen in Figure 5.4. For the second case, since there are more nodes allocated to Task 1 than Task 2, node n_1 scatters its partial applicable exploit list in the local Task 1 communicator, and all other Task 1 nodes follow the same pattern seen in the first case. This second case can be seen in Figure 5.5.

5.3.2.4 Task 2: Each node in Task 2 iterates through the received partial applicable exploit list and creates new states with edges to the current state. However, synchronous firing work is performed during this task, and syncing multiple exploits that could be distributed across multiple nodes leads to additional overhead and complexity. To prevent these difficulties, each node checks its partial applicable exploit list for exploits that are part of a group, removes these exploits from its list, and sends the exploits belonging to a group to the Task 2 local communicator root. Since the Task 2 local root now contains all group exploits, it can execute the Synchronous Firing work without additional communication or synchronization between other MPI nodes in the Task 2 stage. Other than the additional setup steps required for Synchronous Firing for the local root, all work performed during

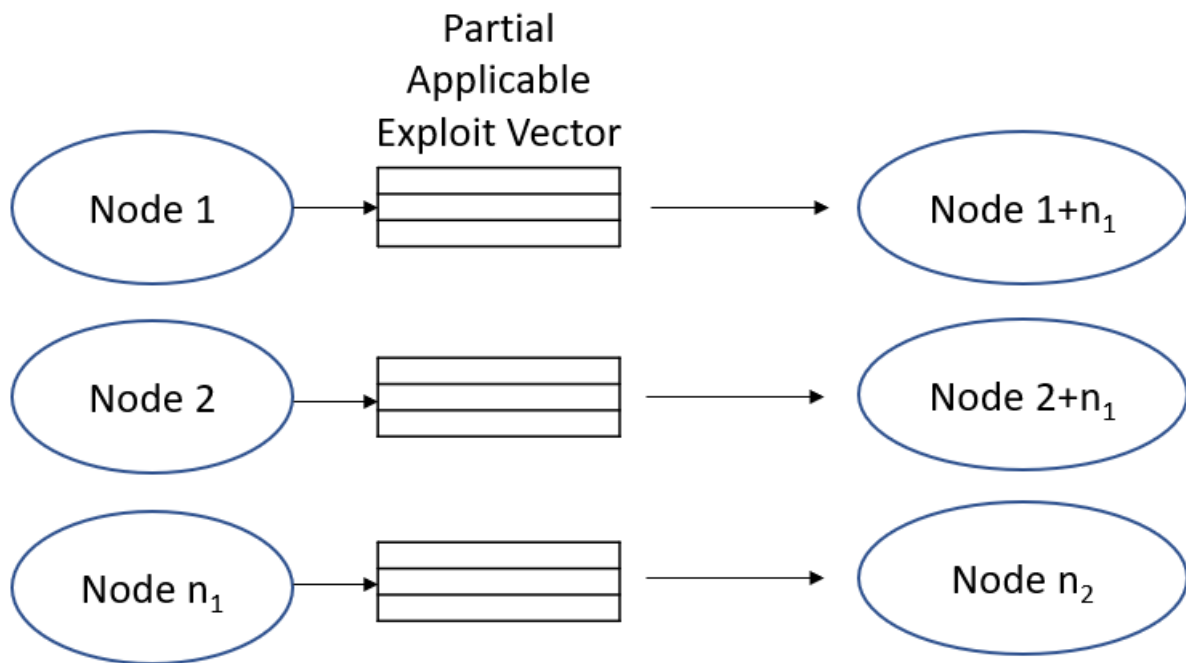


Figure 5.4: Communication From Task 1 to Task 2 when the Number of Nodes Allocated is Equal

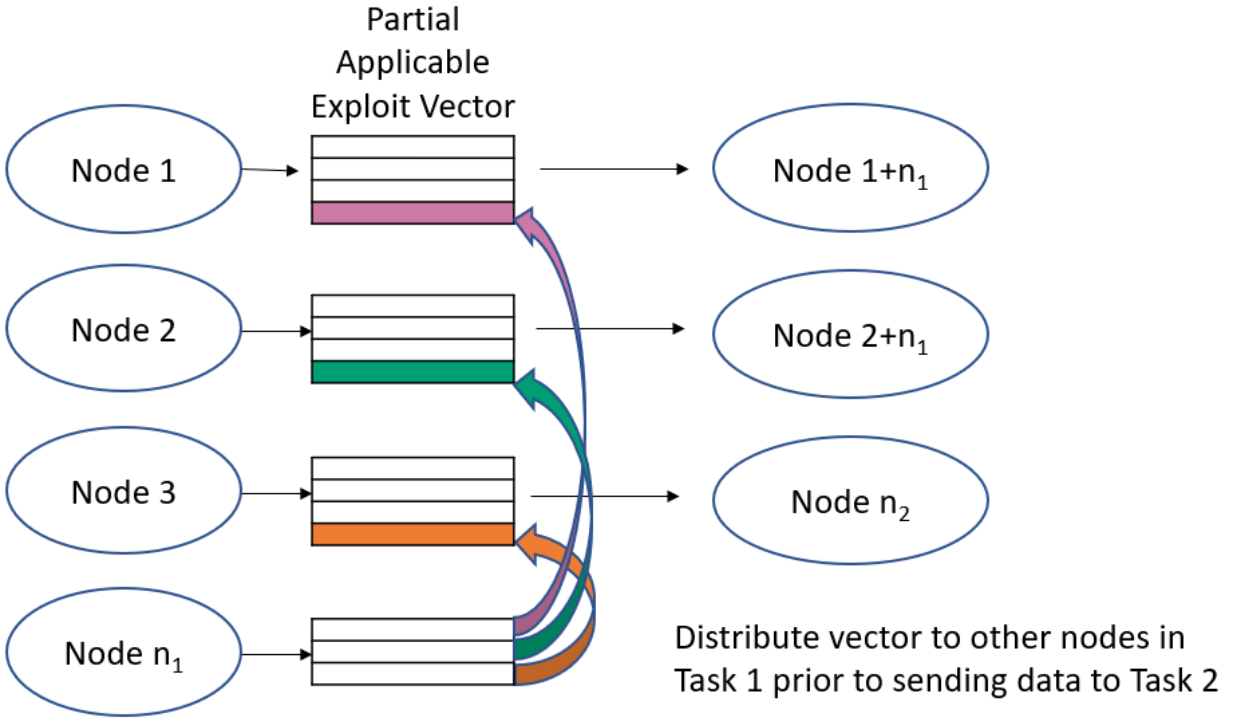


Figure 5.5: Communication From Task 1 to Task 2 when Task 1 Has More Nodes Allocated

this task by all MPI nodes is that seen from the Synchronous Firing figure (Figure 4.3).

5.3.2.5 Task 3: Task 3 is performed only by the root node, and no division of work is necessary. The root node will continuously check for new states until the Task 2 finalize signal is detected. This task consists of setting the new state's ID, adding it to the frontier, adding its information to the instance, and inserting information into the hash map. When the root node has processed all states and has received the Task 2 finalize signal, it will complete Task 3 by sending the instance and/or frontier to Task 4 and/or 5, respectively if applicable, then proceed to Task 0.

5.3.2.6 Task 4 and Task 5: Intermediate database operations, though not frequent and may never occur for small graphs, are lengthy and time-consuming when they do occur. As discussed in Section 3.5, the two main memory consumers are the frontier and the instance, both of which are contained by the root node's memory. Since the database storage

requests are blocking, the pipeline would halt for a lengthy period of time while waiting for the root node to finish potentially two large storages. Tasks 4 and 5 work to alleviate the stall by executing independently of the regular pipeline execution flow. Since Tasks 4 and 5 do not send any data, no other tasks must wait for these tasks to complete. The root node can then asynchronously send the frontier and instance to the appropriate nodes as needed, clear its memory, and continue execution without delay. After initial testing, it was determined that the communication cost of the asynchronous sending of data for Tasks 4 and 5 is less than the time requirement of a database storage operation if performed by the root node.

5.3.2.7 MPI Tags: To ensure that the intended message is received by each node, the MPI message envelopes have their tag fields specified. When a node sends a message, it specifies a tag that corresponds with the data and intent for which it is sent. The tag values were arbitrarily chosen, and tags can be added to the existing list or removed as desired. When receiving a message, a node can specify to only look for messages that have an envelope with a matching tag field. Not only do tags ensure that nodes are receiving the correct messages, they also reduce complexity for program design. Table 5.1 displays the list of tags used for the MPI Tasking approach.

5.3.3 Performance Expectations and Use Cases

Due to the amount of communication between nodes to distribute the necessary data through all stages of the tasking pipeline, this approach is not expected to outperform the serial approach in all cases. This tasking approach was specifically designed to reduce the computation time when the generation of each individual state increases in time. This approach does not offer any guarantees of processing through the frontier at an increased rate; it's main objective is to distribute the workload of individual state generation. As discussed in Section 1.1, the amount of entries in the National Vulnerability database and any custom vulnerability testing to ensure adequate examination of all assets in the network sums to large number of exploits in the exploit list. Likewise for compliance graphs and

Tag	Description
2	Task 2 Finalize Signal
3	Fact for Hash Map Update
4	NetworkState for Hash Map Update
5	NetworkState to be Added to the Frontier
6	Current NetworkState Reference for Edge Creation
7	Factbases for Task 4
8	Edges for Task 4
9	Group Exploit Vectors for Local Root in Task 2
10	Exploit Reference for Task 3 Work
11	AssetGroup Reference for Task 3 Work
14	Continue Signal
15	Finalize Signal
20	Current NetworkState Reference for Task 1
21	Applicable Exploit Vector Scatter for Task 1 Case 2
30	Applicable Exploit Vector Send to Task 2
40	NetworkState Send to Task 2
50	NetworkState to Store in Task 5

Table 5.1: MPI Tags for the MPI Tasking Approach

compliance examinations, Section 1.2.3 discussed that the number of compliance checks for SOX, HIPAA, GDPR, PCI DSS, and/or any other regulatory compliance also sums to a large number of compliance checks in the exploit list. Since the generation of each state is largely dependent on the number of exploits present in the exploit list, this approach is best-suited for when the exploit list grows in size. As will be later discussed, it is also hypothesized that this approach is well-suited when many database operations occur.

5.3.4 Results

A series of tests were conducted on the platform described at the beginning of Section 4.3, and results were collected in regards to the effect of the MPI Tasking approach on increasing sizes of exploit lists for a varying number of nodes. The exploit list initially began with 6 items, and each test scaled the number of exploits by a factor of 2. The final test was with an exploit list with 49,512 entries. If all of the items in these exploit lists were applicable, the runtime would be too great for feasible testing due to the state space explosion. To prevent state-space explosion but still gather valid results, each exploit

list in the tests contained 6 exploits that could be applicable, and all remaining exploits were not applicable. The not applicable exploits were created in a fashion similar to that seen in Figure 5.6. By creating a multitude of not applicable exploits, testing can safely be conducted by ensuring state space explosion would not occur while still observing the effectiveness of the tasking approach.

The results of the Tasking Approach can be seen in Figure 5.7. In terms of speedup, when the number of entries in the exploit list is small, the serial approach has better performance. This is expected due to the communication cost requiring more time than it does to generate a state, as discussed in Section 5.3.3. However, as the number of items in the exploit list increase, the Tasking Approach quickly begins to outperform the serial approach. It is notable that even when the tasking pipeline is not fully saturated (when there are less compute nodes assigned than tasks), the performance is still approximately equal to that of the serial approach. The other noticeable feature is that as more compute nodes are assigned, the speedup continues to increase.

In terms of efficiency, 2 compute nodes offer the greatest value since the speedup using 2 compute nodes is approximately 1.0 as the exploit list size increases. While the 2 compute node option does offer the greatest efficiency, it does not provide a speedup greater than 1.0 on any of the testing cases conducted. The results also demonstrate that an odd number of compute nodes in a fully saturated pipeline has better efficiency than an even number of compute nodes. When referring to Figure 5.2, when there is an odd number number of compute nodes, Task 1 is allocated more nodes than Task 2. In the testing conducted, Task 1 was responsible for iterating through an increased size of the exploit list, so more nodes is advantageous in distributing the workload. However, since many exploits were not applicable, Task 2 had a lower workload where only 6 exploits could be applicable. This will be further elaborated upon in Section 6.2, but it is expected that efficiency will increase for real networks, since nodes in Task 2 will see a realistic workload.

5.4 Subgraphing Approach

```

exploit not_applicable_1(any_car)=
  preconditions:
    quality: any_car, can_fly=true;
  postconditions:
    insert quality: a, flying_car=true;
.

```

Figure 5.6: Example of a Not Applicable Exploit for the MPI Tasking Testing

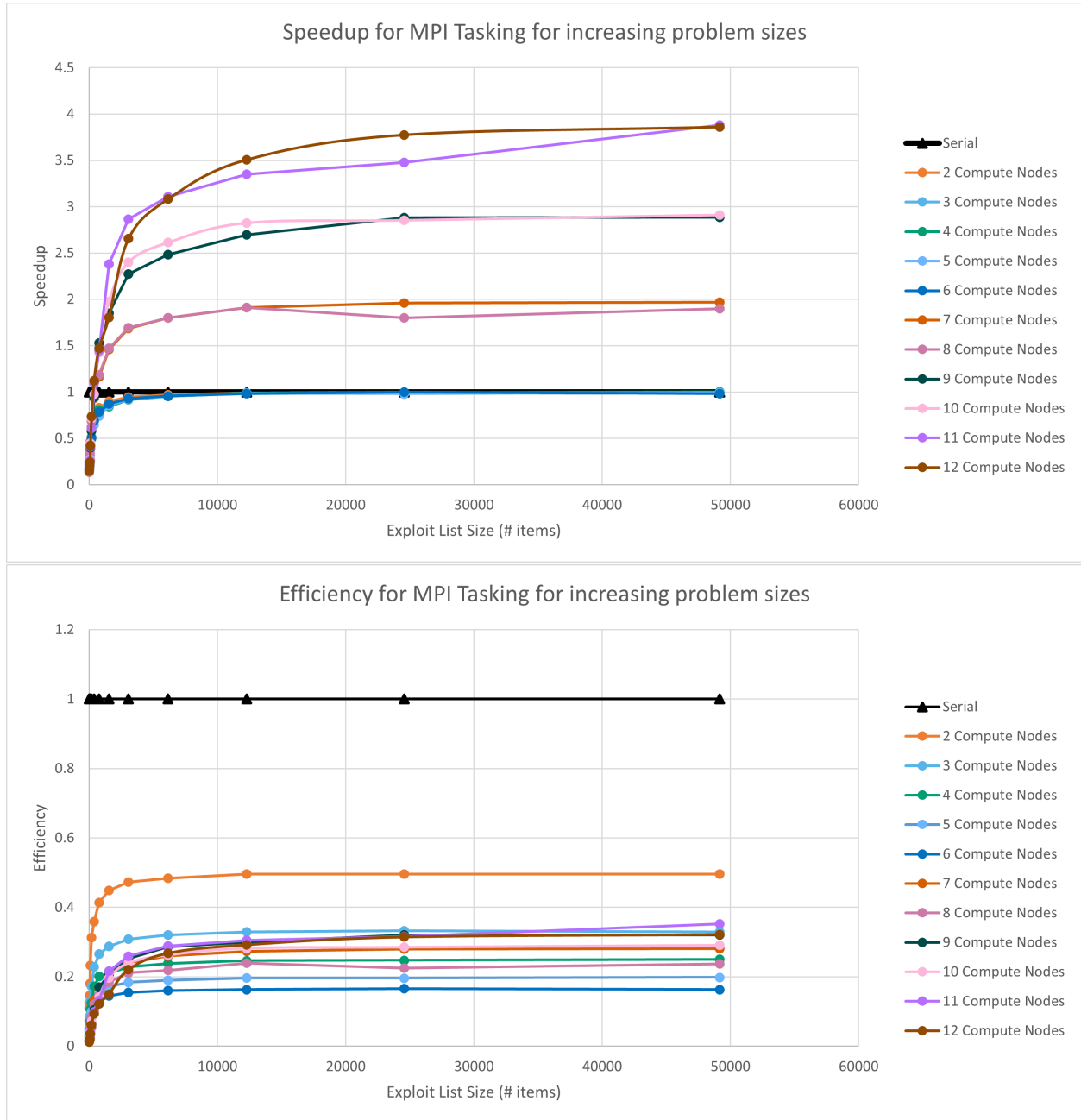


Figure 5.7: Speedup and Efficiency of the MPI Tasking Approach for a Varying Number of Compute Nodes with an Increasing Problem Size

5.4.1 *Introduction to the Subgraphing Approach*

As opposed to the Tasking Approach described in Section 5.3, this Section introduces the Subgraphing Approach as a means of reducing runtime by frontier distribution through subgraphing. Section 3.5 discusses that the frontier is expanded at a rate faster than can be processed. This approach attempts to distribute the frontier by assigning MPI nodes a starting state, and each node will generate a subgraph up to a designated depth-limit, where each node will then return their generated subgraph to a merging process. The authors of [25] presented an alternative method of frontier processing by utilizing OpenMP on a shared-memory system to assign each thread an individual state to explore that would then pass through a critical section. This approach is intended for a distributed system, and additionally differs in that each node will explore multiple states to form a subgraph, rather than exploring one individual state. This approach was implemented with two versions, and both collected results to draw conclusions in regards to speedup, efficiency, and scalability for attack graph and compliance graph generation.

5.4.2 *Algorithm Design*

The design of the subgraphing approach consists of three main components: worker nodes, the root node, and a database node. The original design did not include a database node, but testing through the implementation of the tasking approach discussed in 5.3.2.6 led to the inclusion of a dedicated database node. The overall design is predicated on the root node distributing data to all worker nodes and merging the worker nodes' subgraphs. Figure 5.8 displays an example graph that utilizes three worker nodes with a depth limit of 3. Each worker node corresponds to a different graph state color and texture in the figure. Each worker node explored a varying number of states, but did not proceed to explore a depth that exceeded the specified depth limit of 3. The final cluster of four nodes at the bottom of the graph represents one of the three worker nodes exploring the final states, while the other two nodes wait for further instruction. The following three subsections describe each component in further detail.

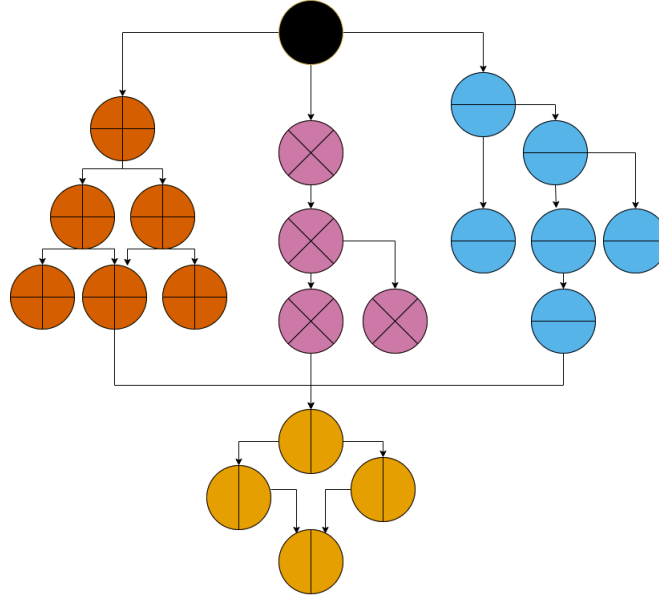


Figure 5.8: Example Graph Using the MPI Subgraphing Approach

5.4.2.1 Worker Nodes: Each worker node will start each iteration of its loop by checking for the finalize message or for a new state to receive. If no message is available to receive, it will continue to wait until one is available. Each worker node will continue this process until the finalize message is received, where it will exit the generation process. When a worker node receives a new state, it will follow the original graph generation process with a modification in the exploration approach. Since other nodes will also be generating subgraphs, using a breadth-first search is unideal since the amount of duplicate work would be increased. Instead, the graph generation will utilize a depth-first search approach, where each node will explore up to a specified depth level. This depth level is specified by the user at before the generation process begins, so additional tuning can be performed as desired. After either the depth limit is reached, or no other states are available to be explored, the node will return its local frontier (if the node still has unexplored states it did not explore) and its generated subgraph. The node will then proceed to wait for a new message from the root node.

5.4.2.2 Root Node: The root node is responsible for two main portions of the subgraphing approach - the distribution of work and the merging of results. The subgraphing approach begins by distributing the only state that is known at the beginning of the generation process to a single node. Once the node returns, two functions may occur. If the node indicates that there are still more states to explore, then the worker node's frontier gets merged with the root node's frontier. If the node discovered new states, then its subgraph gets merged with the root node's graph. This process occurs as long as there are more states to explore. The frontier merging process takes the worker node's frontier and emplaces it to the back of the root node's frontier, with an additional marker. Since each worker node follows a depth-first search, if two worker nodes pop from the front of the queue, there is a high likelihood that the worker nodes are exploring the same segment of the graph, resulting in more duplicate work. To prevent this, each worker node attempts to explore the same segment of the graph throughout the generation process. When the frontiers are merged, the root node includes an additional marker for each worker node to indicate where the end of a node's frontier is in relation to the overall root frontier. When distributing work to a worker node, the root node will pull a state from the queue at a position equal to the node marker. If no marker is found for a worker node and the queue is not empty, then a random state is pulled from the queue instead. Figure 5.9 displays the frontier merging process along with the data distribution.

The root node is also responsible for merging the subgraphs. The root node will receive a subgraph from each node if a subgraph exists, and add all unseen states to the root graph. This is conducted using the hash of each state, since using the state ID is inconsistent across worker nodes. If a state already exists in the root graph, the duplicate state is not added. Likewise, edges are added if the edge has not already been added to the root graph. For the subgraphing approach, edges have been slightly altered from having class members of "From_State_ID" and "To_State_ID" to "From_State_Hash" and "To_State_Hash" to ensure consistency across nodes.

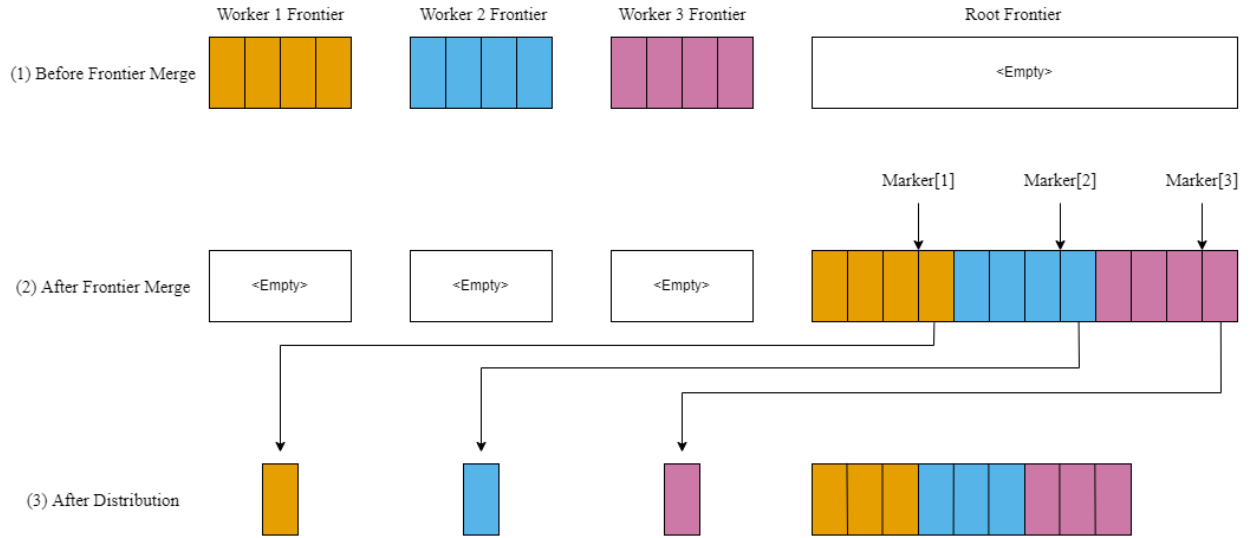


Figure 5.9: Frontier Merging and Data Distribution Process

5.4.2.3 Database Node: A node dedicated to database operations is also present in this approach. The work performed for the subgraphing approach is identical to the work performed by Tasks 4 and 5 in the tasking approach, which is discussed in Section 5.3.2.6.

5.4.2.4 MPI Tags: Similar to Section 5.3.2.7 that discussed the usage of MPI Tags for the tasking approach, the subgraphing approach also made use of MPI Tags. Table 5.2 displays the integers and descriptions of the tags used for this approach.

5.4.3 Performance Expectations and Use Cases

This approach is intended to reduce runtime when the frontier grows at a rate faster than it can be explored. However, since this approach is designed for distributed systems, there is no guarantee that speedup can be achieved when duplicate work is performed. Not only is there wasted time by the worker nodes when duplicate work occurs, but duplicate work also contributes to increased communication between nodes to adequately explore the graph, and also leads to an increased number of merging calls by the root node. The ideal scenario for the subgraphing approach is when the graph is sparse, and the graph aligns more with a N-Ary tree structure where each node only has one parent. When the graph

Tag	Description
1	New State from Root to Workers
2	Workers Finished Subgraphing Signal
3	Local Node Frontiers
5	NetworkState for Hash Map Update
6	Fact for Hash Map Update
7	Factbases for Intermediate Database Storage
8	Edges for Intermediate Database Storage
10	Local Node Factbases
11	Local Node Edges
50	NetworkState for Intermediate Database Storage
99	Generation Finalize Signal

Table 5.2: MPI Tags for the MPI Subgraphing Approach

is sparse, there is a lower likelihood of duplicate work occurring, since worker nodes have a lower probability of exploring a graph state that connects to a different graph state that has been (or will be) explored by another worker node. If each graph state was able to have only one parent, there is a lower likelihood that worker nodes would explore the same graph cluster.

5.4.4 Results

A series of tests were conducted on the platform described at the beginning of Section 4.3, and results were collected in regards to the effect of the MPI Subgraphing approach on the 4 example networks discussed in 4.3.1. All tests used synchronous firing. Figure 5.10 shows the runtimes of each of the test cases. Only the serial runtime from the 2 Service test is displayed for conciseness. The results in terms of speedup and efficiency are seen in Figure 5.11.

As noted from Figures 5.10 and 5.11, the performance from this approach appears quite poor. The serial approach has greater performance in all cases, and the resulting speedups for all 4 service tests are below 1.0 when using the subgraphing approach. Likewise, the efficiency continues to worsen as more compute nodes are added to the system. There are a few reasons for this poor performance. Figure 5.8 illustrates an example graph that is considered favorable to this approach in that branches are relatively distinct, and the

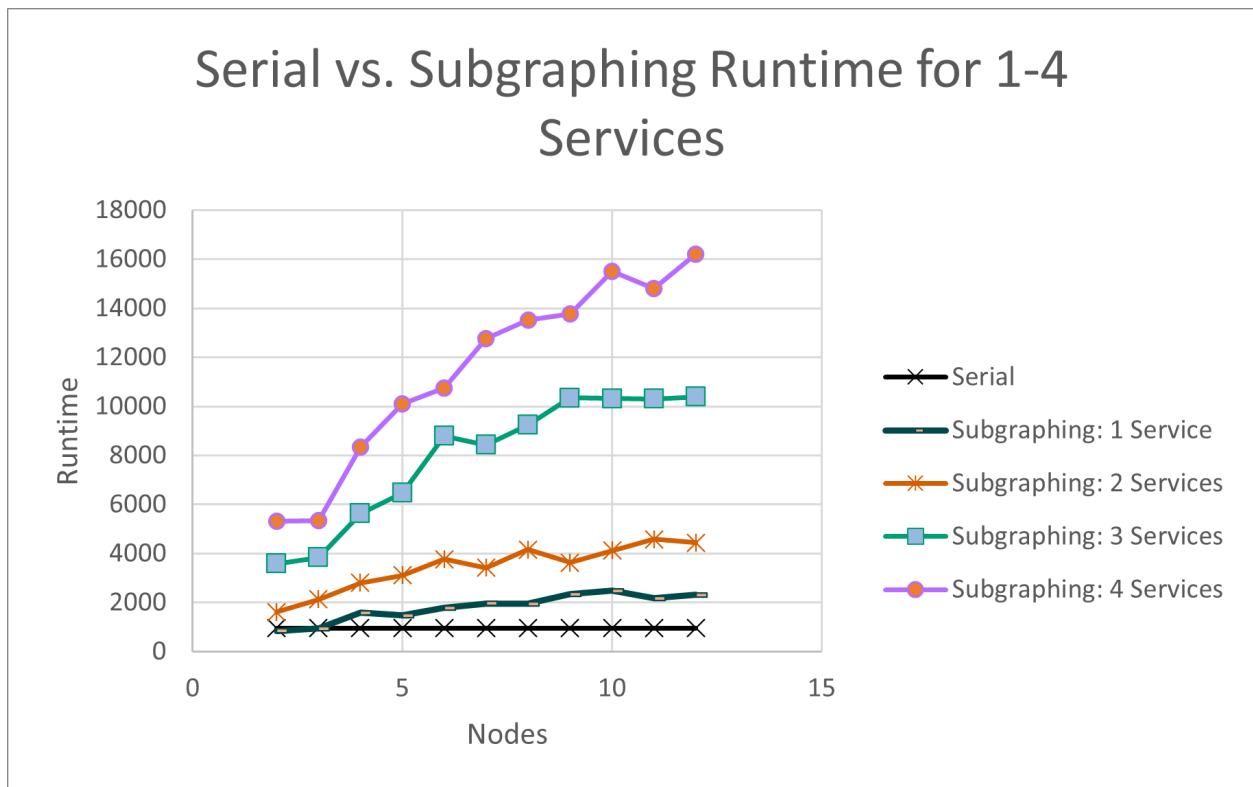


Figure 5.10: First iteration results of MPI Subgraphing in terms of Runtime

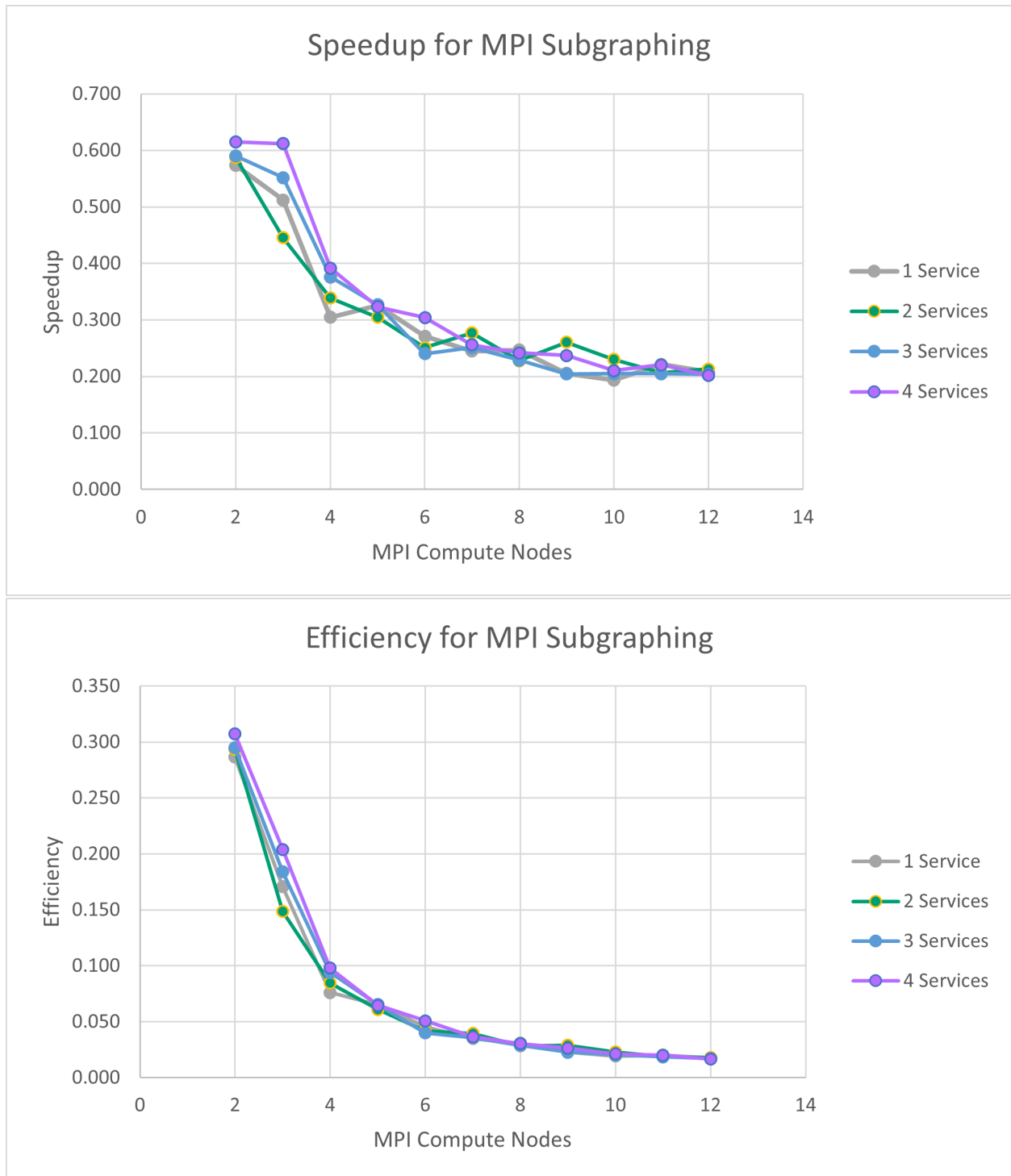


Figure 5.11: First iteration results of MPI Subgraphing in terms of Speedup and Efficiency

graph is not fully connected. As a result in this example graph, each compute node is working on independent graph structures that do not overlap, and all work is distinct. This graph can quickly lead to issues through a few modifications. Figure 5.12 utilizes the same example graph from Figure 5.8, but adds two edges between the outside branches. With this arrangement, the 1st and 3rd compute nodes will perform work that overlaps with the work performed by the 2nd compute node. Both compute node 1 and compute node 3 will explore State 5, and depending on the depth limit, compute nodes 1 and 3 will continue to explore State 5's children, leading to almost all of compute 2's work being duplicated twice. This duplicate work occurs at an alarming rate in the service tests that were performed. Figure 5.13 illustrates that there is an extraordinarily large amount of duplicate work occurring in the testing, which substantially increases the runtime of this approach. The duplicate work, as discussed in Section 5.4.3, not only wastes compute nodes' times and leads to a longer exploration process, but it also requires the root node to perform more merging and cleanup work. When using mpiP (a light-weight MPI profiler provided by Lawrence Livermore National Laboratory) [34], it was measured that this extra merging and cleanup work performed by the root causes additional delays in distributing data, and the compute nodes spent a combined 35% total application runtime just waiting to receive more data from the root node in the 1 service test.

To minimize the duplicate work performed, a second approach using a distributed hash table (DHT) was attempted. With a DHT, each compute node would check to ensure that they were not duplicating work. This would limit the work needed by the root node, but each worker node would need to search the DHT. Using a DHT would increase the communication overhead, but if the communication overhead was less than the time taken for duplicate work or was minimal enough to still process the frontier at a greater rate than the serial approach, then the distributed hash table would be considered advantageous. Rather than devising a unique strategy for a distributed hash table, this work made use of the Berkely Container Library (BCL), which is an open-source library and provides distributed data structures with easy-to-use interfaces. Since BCL is a header-only library, it allowed for

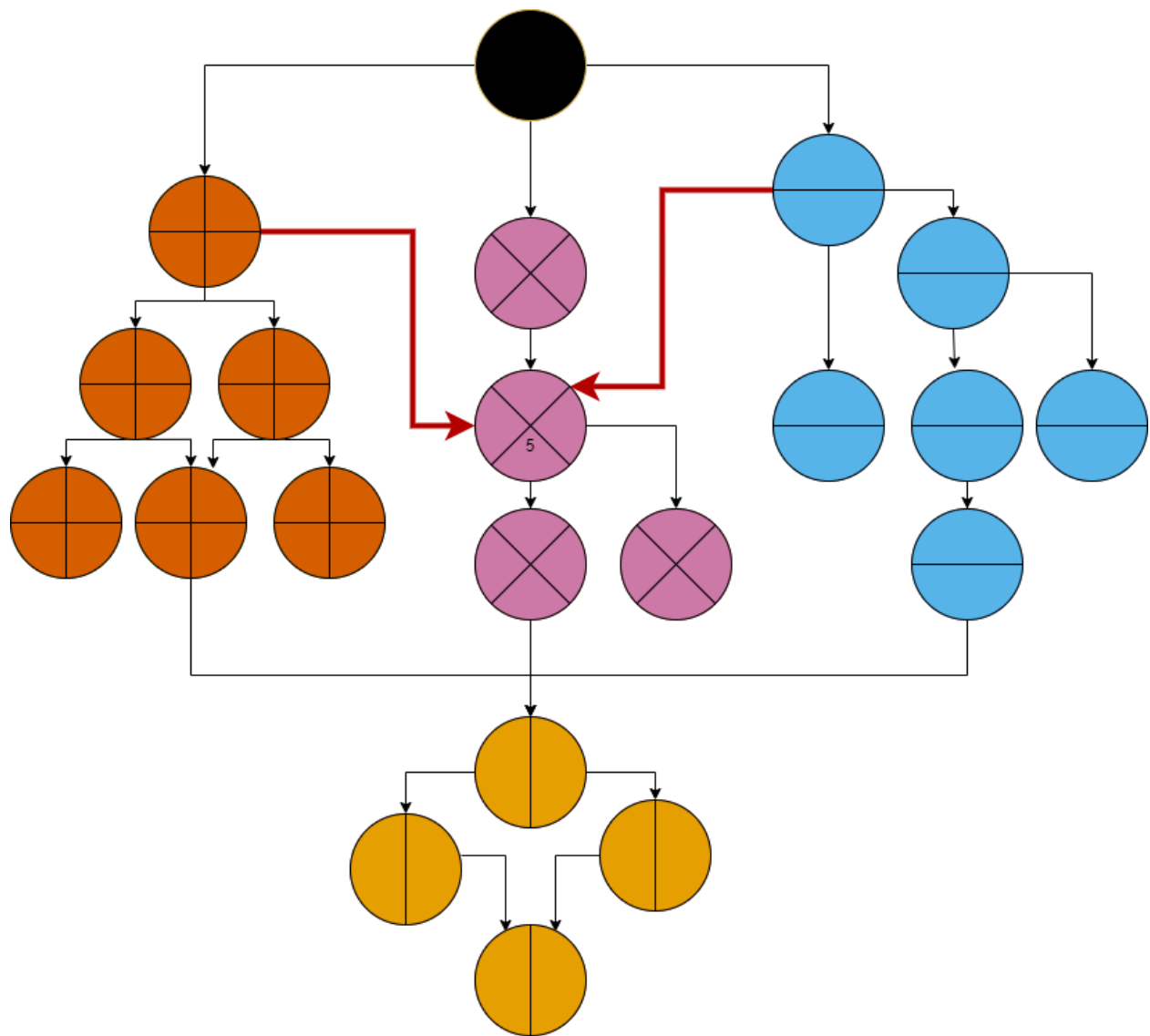


Figure 5.12: Modified Subgraphing Example Graph with Two New Edges

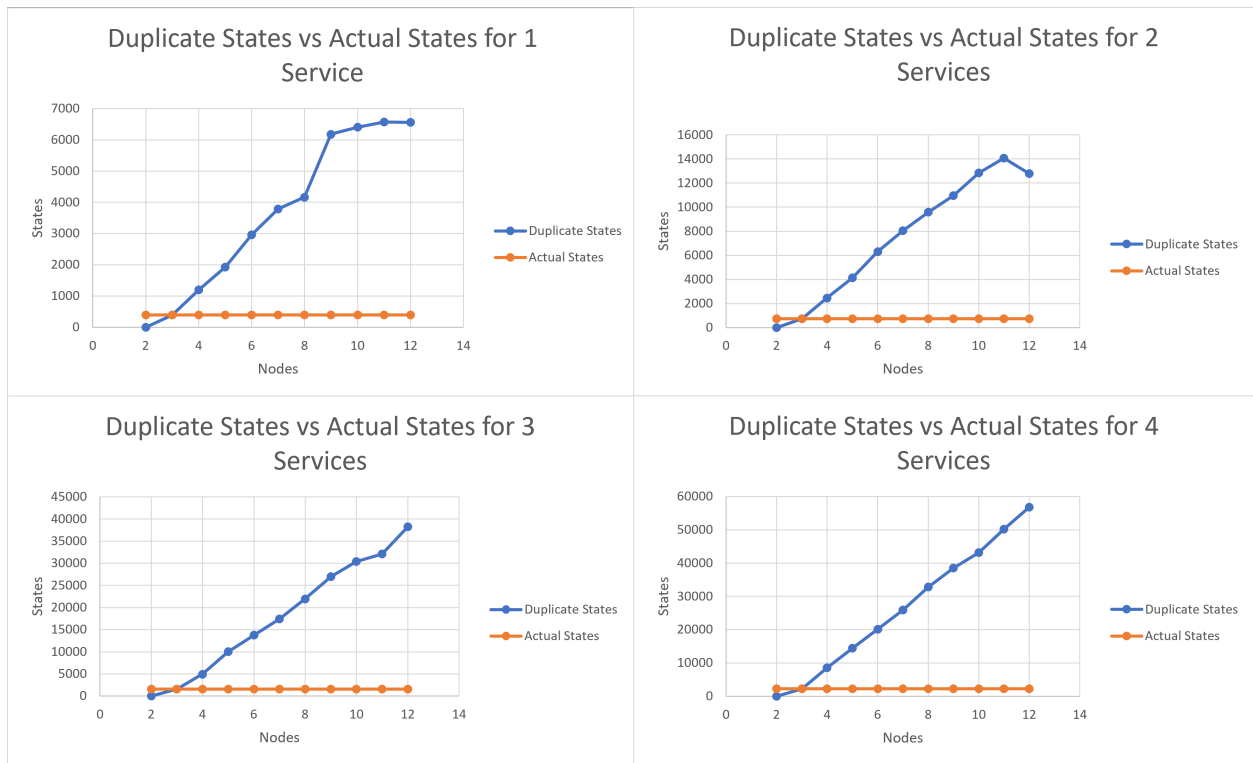


Figure 5.13: Duplicate States Explored vs Actual Number of States for the 1-4 Service Tests

minimal code alterations, and primarily just needed to be dropped into the system. Testing was repeated with an identical setup to the approach without a DHT. The results in terms of speedup and efficiency are seen in Figure 5.14. Results in terms of runtime between the DHT approach and the base approach are seen in Figure 5.15.

Implementing the DHT did prevent duplicate work, but the communication cost from repeated DHT queries by each worker node was far greater than the serial approach, and was also greater than the first approach for MPI Subgraphing without the DHT. As a result, the MPI Subgraphing approach is not viable as it stands. Future improvements or entire reworking will need to be performed, and this is discussed further in Section 6.2.

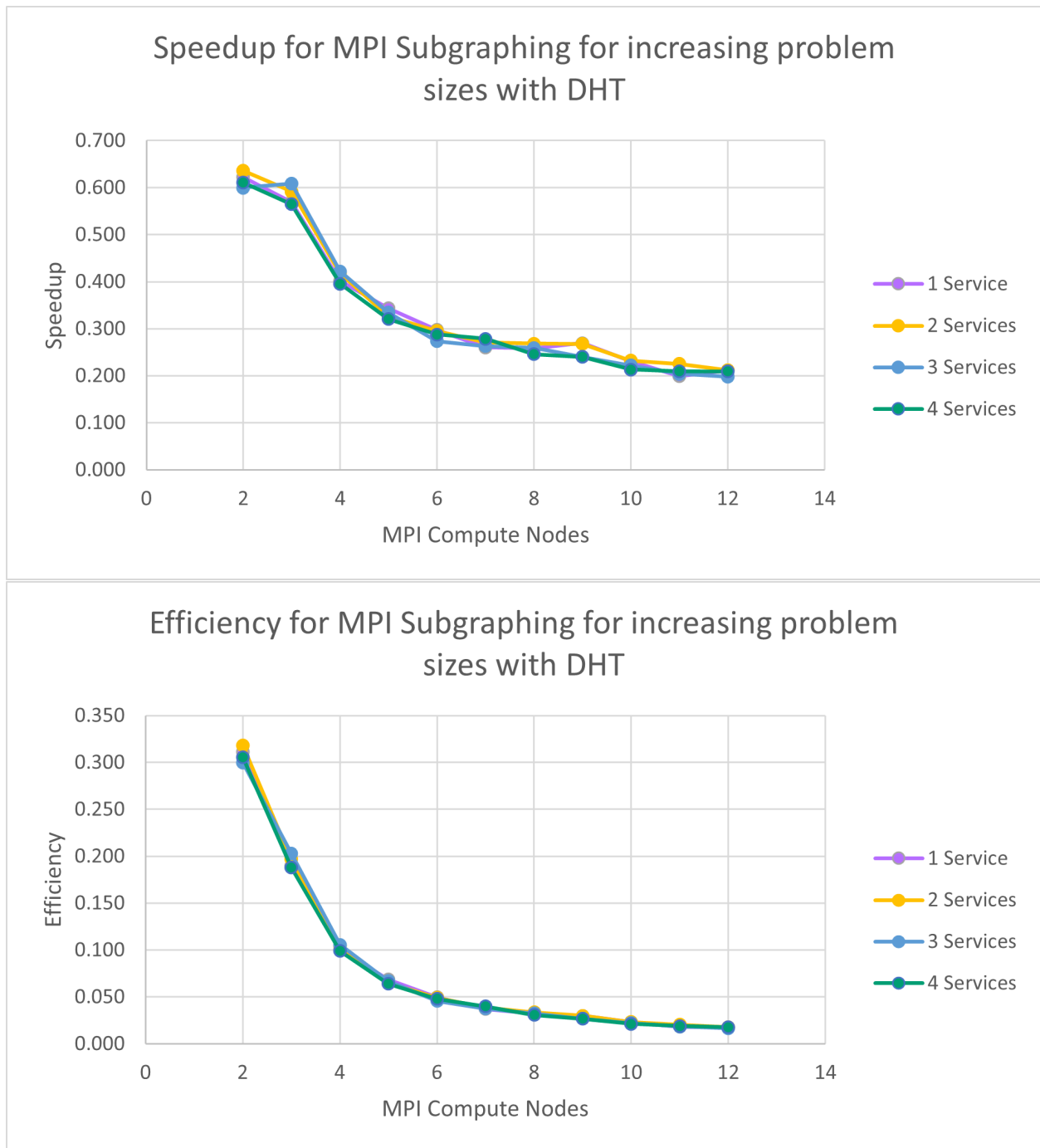


Figure 5.14: Speedup and Efficiency of MPI Subgraphing when using a DHT

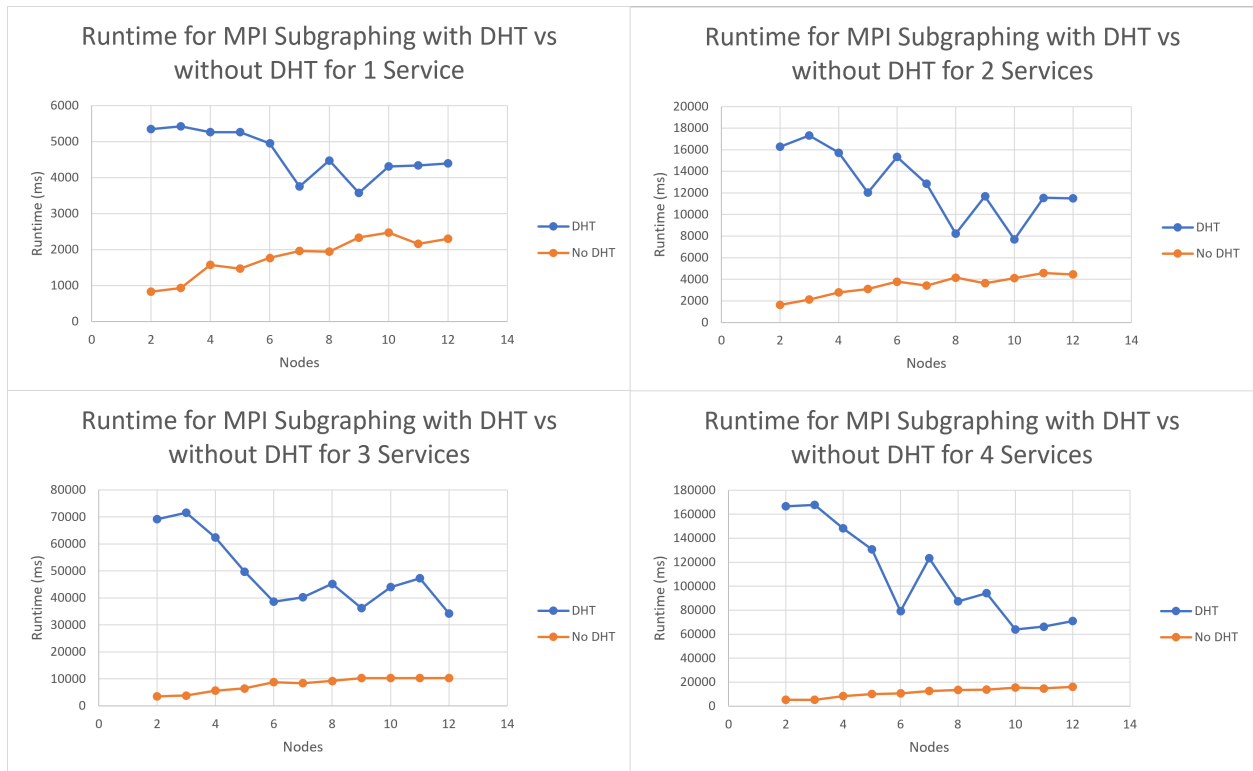


Figure 5.15: Runtime of MPI Subgraphing when using a DHT vs not using a DHT

CHAPTER 6

CONCLUSIONS AND FUTURE WORKS

6.1 Conclusions

This thesis presented various extensions to an attack graph generator, RAGE, to allow for a broader range of utilities. In order to reduce the complexity required for network model and exploit file creations, Sections 3.3 and 3.4 discussed the implementation of relational and compound operators. Both implementations simplify the amount of manual enumeration and manual specifications of asset qualities, and this simplifies precondition checks to singular lines. In addition, these implementations reduce the complexity required for synchronous firing exploit creations by avoiding the need for (for instance) time flags and enumeration of time all time instances. These expansions also allow for more complex attack modeling, since broad sweeps and generic $<$ or $>$ checks can be performed. Due to the intermediate database storage feature presented in Section 3.5, very large attack or compliance graphs can be generated without concern of absolute memory consumption, assuming very large storage solutions are in place. Section 3.1 discusses the path walking feature, which is able to split attack and compliance graphs into subgraphs that can be used to simplify the analysis process by examining only smaller portions or focus areas of a network at a time.

Chapter 4 presents the synchronous firing feature, which is successfully able to reduce the state space and runtime of the generation process when assets have inseparable features. This feature does not lose any substantive information from a network and its resulting graph; the graph is able to remain exhaustive and still capture all necessary information. The results are promising, and greater reductions are expected when a greater number of assets share inseparable features, as discussed in Section 6.2.

Chapter 5 presented two approaches for utilizing MPI for extension to the distributed

computing platform space. One approach was a task parallelism approach discussed in Section 5.3, and promising results were observed when the generation of each state increased in computation requirements. The second approach was a data parallelism approach discussed in Section 5.4. While results were not promising for this approach, future work can be conducted to optimize and avoid the difficulties of duplicate work and communication overhead. This thesis provides an MPI code-base for parallelizing the generation of compliance graphs on multiple nodes.

Throughout this thesis and its works, RAGE has demonstrated its extensions to support compliance graph generation. Section 1.2.2 discussed the alterations required for attack graph generators to support compliance graphs, and example compliance graphs have been generated in the results seen in Sections 4.3.2, 5.3.4, and 5.4.4.

6.2 Future Work

There are multiple avenues that future works and research can be conducted. One such investigation involves examining the effect of the synchronous firing feature with more assets belonging to groups. As the number of assets with inseparable features increases, more permutations consisting of unattainable states in a graph would exist. Since more assets can devolve and grow increasingly out of sync, the reduction rate when using the synchronous firing feature is likely to increase.

Another avenue is to improve and further test the MPI Tasking approach discussed in Section 5.3. One improvement approach involves blending OpenMP and MPI. For Task 2 specifically, an example diagram was created that displays a possible approach at leveraging both OpenMP and MPI, and this can be seen in Figure 6.1. Similar techniques could likely be implemented for Task 1, and other techniques could be devised for these tasks and for Tasks 3, 4, and 5. Further testing can be conducted to examine the effectiveness of MPI Tasking when generating large graphs with many applicable exploits at any given state. This would allow Task 2 to have a greater workload, where increased speedups are likely to be attained over the serial approach.

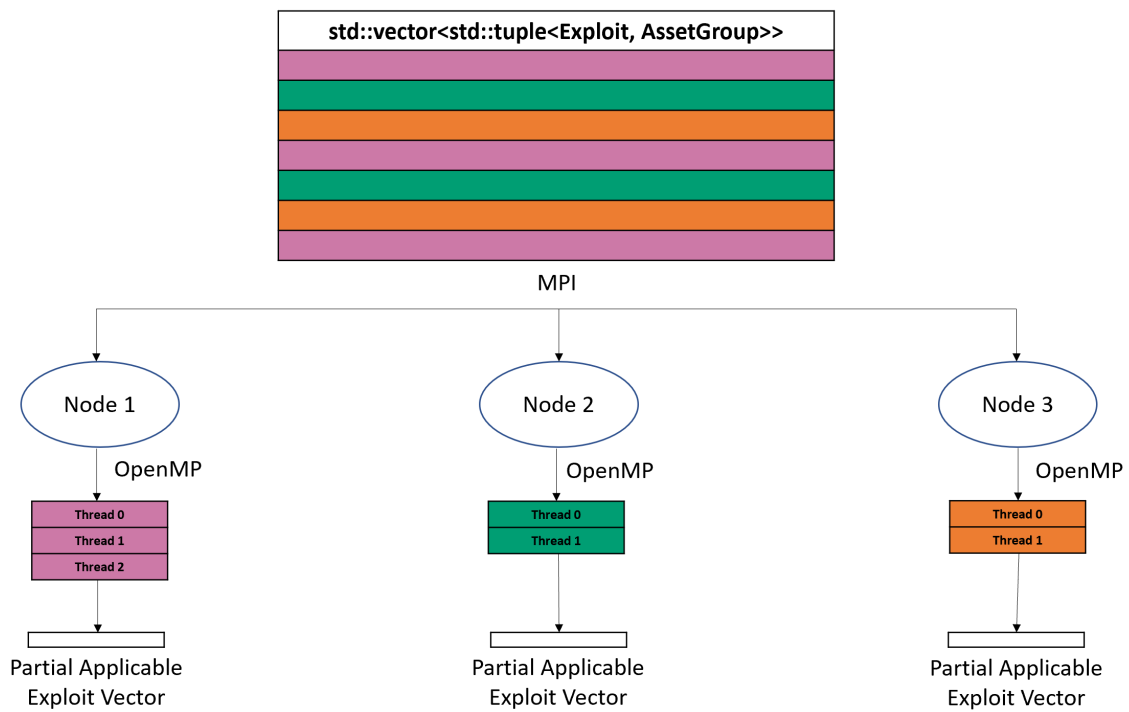


Figure 6.1: Possible Method for Blending MPI and OpenMP for Task 2 of the MPI Tasking Approach

Many research avenues exist for the MPI Subgraphing approach seen in Section 5.4. The results collected in this work were not promising, and many alterations can be attempted. Duplicate states are the primary problem point of this approach, and using a DHT increased the communication cost to a large degree. Alternate methods for addressing this issue can be devised. While it did not improve the results in this work, parameter tuning with this approach can be performed in regards to the depth limit value and the metric for when a worker node does not have a queue marker. This work used a random metric, but more intelligent strategies could be devised.

Drastic results are not likely to occur, but various MPI parameter tunings can be performed for both the MPI Tasking and Subgraphing approaches. For instance, this work utilized the parameters for binding to and mapping by NUMA. This is seemingly the most common approach, and MPI balances the workload of NUMA nodes automatically when these parameters are used. Many other options exist for these parameters, and many other parameters exist that can be examined. In addition, comparisons between OpenMPI, MPICH, and MVAPICH can be conducted to see if there is any noticeable, consistent improvement.

BIBLIOGRAPHY

- [1] C. Phillips and L. P. Swiler, “A graph-based system for network-vulnerability analysis,” *Proceedings New Security Paradigms Workshop*, vol. Part F1292, pp. 71–79, 1998. doi: 10.1145/310889.310919.
- [2] B. Schneier, “Modeling Security Threats,” *Dr. Dobb’s Journal*, 1999. vol. 24, no.12.
- [3] X. Ou, W. F. Boyer, and M. A. Mcqueen, “A Scalable Approach to Attack Graph Generation,” *CCS ’06: Proceedings of the 13th ACM conference on Computer and communications security*, pp. 336–345, 2006.
- [4] O. Sheyner, J. Haines, S. Jha, R. Lippmann, and J. Wing, “Automated Generation and Analysis of Attack Graphs,” *Proceeding of 2002 IEEE Symposium on Security and Privacy*, pp. 254–265, 2002.
- [5] J. Zhang, S. Khoram, and J. Li, “Boosting the performance of FPGA-based graph processor using hybrid memory cube: A case for breadth first search,” *FPGA 2017 - Proceedings of the 2017 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays*, pp. 207–216, 2017.
- [6] J. Hale, P. Hawrylak, and M. Papa, “Compliance Method for a Cyber-Physical System.” U.S. Patent Number 9,471,789, Oct. 18, 2016.
- [7] N. Baloyi and P. Kotzé, “Guidelines for Data Privacy Compliance: A Focus on Cyber-physical Systems and Internet of Things,” in *SAICSIT ’19: Proceedings of the South African Institute of Computer Scientists and Information Technologists 2019*, (Skukuza South Africa), Association for Computing Machinery, 2019.

- [8] E. Allman, “Complying with Compliance: Blowing it off is not an option,” *ACM Queue*, vol. 4, no. 7, 2006.
- [9] “Sarbanes-Oxley Act of 2002.” Pub. L. No. 107-204. 2002 [Online]. Available: <https://www.govinfo.gov/content/pkg/PLAW-107publ204/html/PLAW-107publ204.htm>.
- [10] “Health Insurance Portability and Accountability Act of 1996.” Pub. L. No. 104-191. 1996 [Online]. Available: <https://www.govinfo.gov/content/pkg/PLAW-104publ191/html/PLAW-104publ191.htm>.
- [11] T. E. Parliment and the Council of the European Union, “Regulation (eu) 2016/679 of the European Parliament and of the Council of 27.” Available: <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32016R0679>.
- [12] P. S. S. Council, “Payment Card Industry (PCI) Data Security Standard,” May 2018. Available: https://www.pcisecuritystandards.org/documents/PCI_DSS_v3-2-1.pdf.
- [13] K. Cook, *RAGE: The Rage Attack Graph Engine*. PhD thesis, The University of Tulsa, 2018.
- [14] J. Berry and B. Hendrickson, “Graph Analysis with High Performance Computing,” *Computing in Science and Engineering*, 2007.
- [15] S. Ainsworth and T. M. Jones, “Graph prefetching using data structure knowledge,” *Proceedings of the International Conference on Supercomputing*, vol. 01-03-June, 2016.
- [16] P. Yao, L. Zheng, X. Liao, H. Jin, and B. He, “An efficient graph accelerator with parallel data conflict management,” *Parallel Architectures and Compilation Techniques - Conference Proceedings, PACT*, 2018.
- [17] G. Dai, Y. Chi, Y. Wang, and H. Yang, “FPGP: Graph processing framework on FPGA: A case study of breadth-first search,” *FPGA 2016 - Proceedings of the 2016*

- ACM/SIGDA International Symposium on Field-Programmable Gate Arrays*, pp. 105–110, 2016.
- [18] S. Arifuzzaman and M. Khan, “Fast parallel conversion of edge list to adjacency list for large-scale graphs,” in *HPC ’15: Proceedings of the Symposium on High Performance Computing*, pp. 17–24, Apr. 2015.
 - [19] X. Yu, W. Chen, J. Miao, J. Chen, H. Mao, Q. Luo, and L. Gu, “The Construction of Large Graph Data Structures in a Scalable Distributed Message System,” in *HPCC 2018: Proceedings of the 2018 2nd High Performance Computing and Cluster Technologies Conference*, pp. 6–10, June 2018.
 - [20] P. Liakos, K. Papakonstantinou, and A. Delis, “Memory-Optimized Distributed Graph Processing through Novel Compression Techniques,” in *CIKM ’16: Proceedings of the 25th ACM International Conference on Information and Knowledge Management*, pp. 2317–2322, Oct. 2016.
 - [21] J. Balaji and R. Sunderraman, “Graph Topology Abstraction for Distributed Path Queries,” in *HPGP ’16: Proceedings of the ACM Workshop on High Performance Graph Processing*, pp. 27–34, May 2016.
 - [22] “An Overview of the Parallel Boost Graph Library - 1.75.0,” 2009.
 - [23] J. Siek, L.-Q. Lee, and A. Lumsdaine, “The Boost Graph Library, vers. 1.75.0.” https://www.boost.org/doc/libs/1_75_0/libs/graph/doc/index.html.
 - [24] K. Cook, T. Shaw, J. Hale, and P. Hawrylak, “Scalable attack graph generation,” *Proceedings of the 11th Annual Cyber and Information Security Research Conference, CISRC 2016*, 2016.
 - [25] M. Li, P. Hawrylak, and J. Hale, “Concurrency Strategies for Attack Graph Generation,” *Proceedings - 2019 2nd International Conference on Data Intelligence and Security, ICDIS 2019*, pp. 174–179, 2019.

- [26] M. Li, P. J. Hawrylak, and J. Hale, “Implementing an attack graph generator in cuda,” in *2020 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW)*, pp. 730–738, 2020.
- [27] K. Kaynar and F. Sivrikaya, “Distributed attack graph generation,” *IEEE Transactions on Dependable and Secure Computing*, vol. 13, no. 5, pp. 519–532, 2016.
- [28] The Graphviz Authors, “DOT Language.” <https://graphviz.org/doc/info/lang.html>.
- [29] W. M. Nichols, *Hybrid Attack Graphs for Use with a Simulation of a Cyber-Physical System*. PhD thesis, The University of Tulsa, 2018.
- [30] “set-value is vulnerable to Prototype Pollution in versions lower than 3.0.1. The function mixin-deep could be tricked into adding or modifying properties of Object.prototype using any of the constructor, prototype and _proto_ payloads..” National Vulnerability Database, Aug. 2019.
- [31] M. Li, P. Hawrylak, and J. Hale, “Combining OpenCL and MPI to support heterogeneous computing on a cluster,” *ACM International Conference Proceeding Series*, 2019.
- [32] G. Louthan, *Hybrid Attack Graphs for Modeling Cyber-Physical Systems*. PhD thesis, The University of Tulsa, 2011.
- [33] P. Pacheco, *An Introduction to Parallel Programming*. Morgan Kaufmann, print ed., 2011.
- [34] Lawrence Livermore National Laboratory, “mpiP, a light-weight MPI profiler.” <https://software.llnl.gov/mpiP/>.