

Article

Not peer-reviewed version

Effects of Poor Workload Partitioning on System Performance for Chiplet-Based Systems

[Peter Mbua](#)*, Forcha Peter, [Christophe Bobda](#)*

Posted Date: 6 February 2026

doi: 10.20944/preprints202602.0486.v1

Keywords: chiplet-based systems; workload partitioning; task mapping; inter-chiplet communication; communication latency; network congestion



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

Effects of Poor Workload Partitioning on System Performance for Chiplet-Based Systems

Peter Mbua , Forcha Peter and Christophe Bobda *

University of Florida

* Correspondence: cbobda@ece.ufl.edu; Tel.: (352) 294-2024

Abstract

The emergence of chiplet-based architectures represents a paradigm shift in post-Moore's Law computing systems, offering substantial cost and yield advantages through functional disaggregation. However, the heterogeneity of inter-chiplet communication introduces unique performance challenges that conventional partitioning strategies fail to address. This work presents a comprehensive characterization of how poor workload partitioning degrades communication performance in chiplet-based systems. We demonstrate, through detailed experimental analysis, that suboptimal workload partitioning can increase inter-chiplet communication latency by up to 10× and can inflate network congestion beyond sustainable levels as systems scale. Our findings show that optimized partitioning strategies can achieve 87.4% reduction in inter-chiplet traffic, improve system throughput by 8.75×, and enhance energy efficiency by 10.3× compared to naive partitioning approaches. We further characterize how these effects compound with system scalability, revealing that communication overhead can consume 85% of execution time in poorly partitioned 16-chiplet systems, versus only 35% in well partitioned configurations. This work provides essential insights into the communication-aware design space of chiplet systems and validates the critical importance of sophisticated workload partitioning algorithms.

Keywords: chiplet-based systems; workload partitioning; task mapping; inter-chiplet communication; communication latency; network congestion

1. Introduction

The semiconductor industry faces fundamental challenges in sustaining Moore's Law while managing escalating design complexity and manufacturing costs. The transition to advanced technology nodes has rendered monolithic chip design economically and physically infeasible for many applications, particularly in high-performance computing and artificial intelligence domains. Chiplet-based architectures, which disaggregate complex systems into smaller functional units manufactured separately and then integrated through advanced packaging technologies, have emerged as a compelling solution [1]. This communication asymmetry fundamentally changes the optimization problem: naive workload distribution strategies that ignore communication patterns inevitably lead to excessive inter-chiplet traffic, congestion, and performance degradation.

Despite this critical challenge, many existing workload partitioning approaches treat all chiplets as identical and employ uniform distribution strategies [2]. Such approaches fundamentally overlook the heterogeneity of the communication substrate and fail to account for how specific workload communication patterns interact with the underlying interconnect architecture. The consequences are substantial: research indicates that chiplet-based systems can forfeit over one-third of monolithic performance due to communication inefficiencies [3], much of which stems from poor workload partitioning decisions.

This work addresses this gap by comprehensively characterizing how the quality of workload partitioning directly influences communication performance in chiplet systems. We systematically quantify the relationship between partitioning strategies and critical performance metrics, including

inter-chiplet communication latency, network congestion, data locality, and overall system throughput. Our experimental analysis demonstrates that the effects of poor partitioning are non-linear and scale with the system, creating increasingly severe bottlenecks as the number of chiplets increases.

The rest of this work is divided as follows: Section 2 discusses important foundational chiplet concepts and highlights the research gap; Section 3 provides details of the experimental study, modeling, and simulation; Section 4 and 5 present the experimental results and discuss trends respectively. Some important ongoing work in optimizing chiplet architectures was discussed in Section 6, and Section 7 summarizes the outcome of this work.

2. Background and Related Work

2.1. Chiplet-Based System Architecture

Chiplet-based systems represent a fundamental departure from traditional monolithic system-on-chip (SoC) designs. In this paradigm, complex functionality is partitioned across multiple independent dies, each manufactured using optimal processes for its specific function, and then integrated through advanced packaging techniques such as 2.5D silicon interposer integration or 3D stacking with through-silicon vias (TSVs) [4].

For deep neural network acceleration and high-performance computing workloads, chiplet-based architectures have demonstrated compelling performance, cost, and power trade-offs. The Simba system, a 36-chiplet prototype for deep learning inference, exemplifies this approach, achieving 4 TOPS per chiplet, with package-level performance of 128 TOPS and an energy efficiency of 6.1 TOPS/W [5]. Similarly, systems like NN-Baton provide hierarchical frameworks for DNN workload orchestration across multiple computation levels in chiplet hierarchies, enabling 22.5%-44% energy savings compared to monolithic alternatives under equivalent configurations [6]. Figure 2.1 shows a typical modern application of system-based integration.

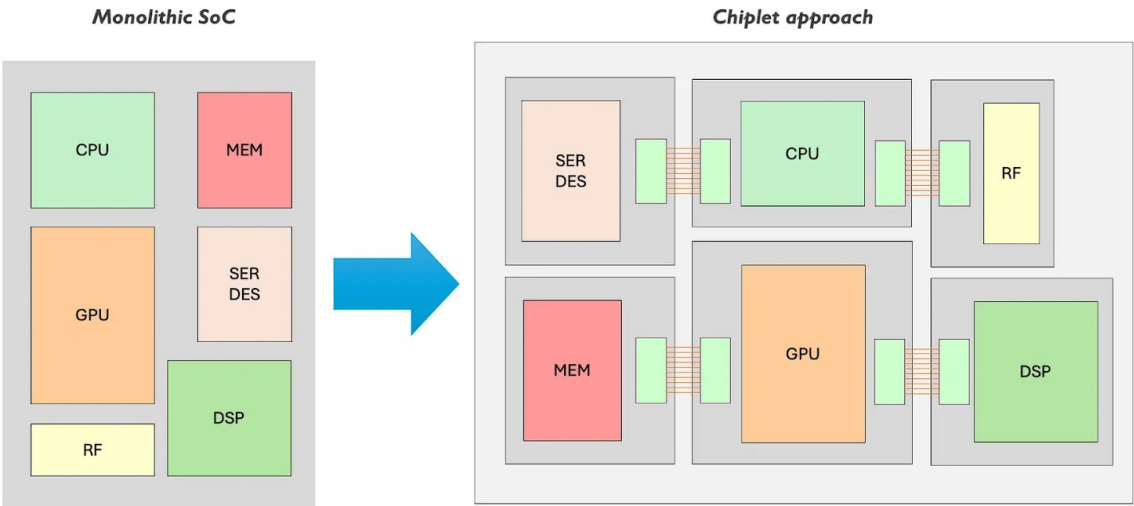


Figure 1. Chiplet heterogeneous integration.

2.2. Communication Challenges in Multi-Chiplet Systems

Despite their advantages, chiplet-based architectures introduce fundamental communication challenges that fundamentally constrain performance. Multi-chiplet systems are naturally Non-Uniform Memory Access (NUMA) systems that suffer from slow remote accesses, with limited throughput and higher inter-chiplet communication latency relative to traditional systems [2]. These communication overheads stem from multiple sources: the physical distance between chiplets on the interposer, routing through limited inter-chiplet pathways, and the serialization bottlenecks inherent in package-level interconnects.

The performance implications are substantial. Research demonstrates that while chiplet-based chips can reduce manufacturing costs by approximately 50%, they simultaneously incur performance

penalties exceeding one-third compared to monolithic designs when communication patterns are not carefully optimized [2]. This performance-cost tradeoff motivates sophisticated workload partitioning strategies that minimize inter-chiplet communication.

Also, system partitioning and architecture definition constitute the foundational stage in the chiplet-based design flow (see Figure 2), directly shaping how computation, memory, and communication workloads are decomposed and mapped onto heterogeneous chiplets. At this stage, the monolithic system specification is transformed into a set of functionally coherent partitions, each representing a candidate chiplet or tightly coupled cluster of chiplets. The primary objective is to determine which functionality belongs where before making assumptions about die-to-die (D2D) interfaces, physical layout, or packaging constraints [7].

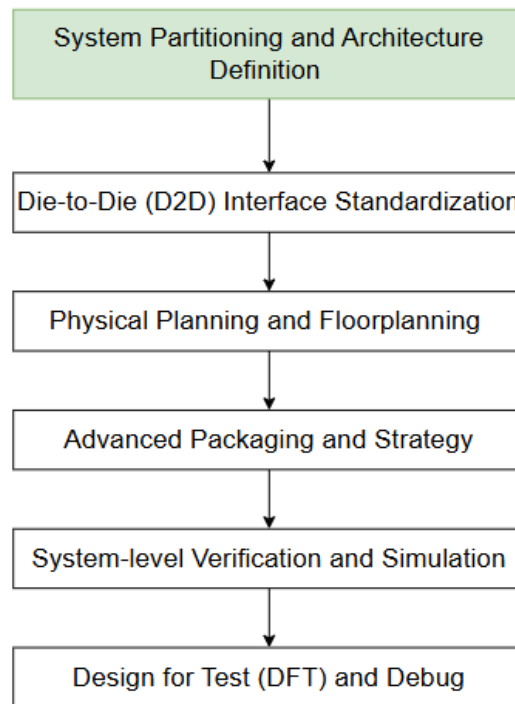


Figure 2. Chiplet-based Design Flow.

In chiplet-based systems, effective partitioning is fundamentally a workload-aware exercise. Computational kernels, data access patterns, control paths, and real-time constraints must be analyzed jointly to ensure that each partition aligns with its performance, power, and scalability requirements. For example, data-intensive and latency-sensitive workloads (e.g., near-sensor processing or packet parsing) are often isolated into dedicated accelerators, while control-dominated or infrequently used functions are mapped to general-purpose or low-power chiplets. Poor early partitioning can lead to excessive inter-chiplet communication, bandwidth bottlenecks, and inefficient use of advanced packaging technologies, all of which are difficult to correct later in the flow [8].

2.3. Task Mapping and Workload Partitioning

The problem of optimally mapping tasks and workloads onto chiplet-based systems has received increasing research attention. Recent work on task mapping in multi-chiplet-based many-core systems reveals that traditional mapping algorithms fail to account for the latency and bandwidth differences between intra-chiplet and inter-chiplet communications, leading to sub-optimal performance [9]. This work proposes a two-step approach combining binary linear programming for task-to-chiplet assignment with latency-minimizing intra-chiplet mapping. The results are impressive: compared to existing methods, the proposed algorithm achieves 37.5%-24.7% reductions in execution time and up to 43.2%-32.9% reductions in communication latency.

Similarly, the VariPar framework demonstrates that accounting for real-world performance variations across chiplets, including manufacturing process variation, thermal conditions, physical placement, and power supply variations, enables significant performance improvements [3]. By modeling performance variations for each chiplet and partitioning workloads accordingly, VariPar achieves 1.45× performance and 1.82× energy-efficiency improvements over naive uniform partitioning strategies. This work conclusively demonstrates that the chiplet-aware partitioning problem is fundamentally different from traditional homogeneous system partitioning.

2.4. Interconnect Network Design and Dataflow Mapping

Beyond simple task mapping, the interaction between workload partitioning and interconnect network design is critical. INDM (Chiplet-Based Interconnect Network and Dataflow Mapping) proposes co-optimizing the interconnect topology and the dataflow mapping strategy for DNN accelerators [10]. By proposing hierarchical interconnect networks with multiring on-die networks and cluster-based inter-die networks, along with communication-aware dataflow mapping to minimize traffic congestion, INDM achieves 26.00%-73.81% energy-delay-product reduction and 26.93%-79.78% latency reduction compared to state-of-the-art approaches.

The M2M framework extends this to multiple simultaneous DNNs on multi-chiplet architectures, proposing temporal and spatial task scheduling for reconfigurable dataflow accelerators and communication-aware task mapping [11]. The framework includes fine-tuned quality-of-service policies for network-on-package links, achieving latency reductions of 7.18%-61.09% across diverse vision, language, and mixed workloads.

2.5. Problem Statement and Research Gaps

While substantial research addresses workload partitioning and task mapping in chiplet systems, a comprehensive characterization of how poor partitioning degrades communication performance remains lacking. Specifically, existing work primarily focuses on demonstrating improvements achieved through optimized partitioning strategies, rather than systematically characterizing the degradation caused by naive approaches. Furthermore, the non-linear scaling effects of poor partitioning as system size increases, which is a critical concern for future systems, have not been thoroughly characterized.

In this work, partitioning refers to the process of deciding which computational tasks run on which chiplet in a multi-chiplet system. Thus, this work addresses these gaps by systematically measuring communication performance metrics across a spectrum of partitioning quality levels, from completely random allocation to optimal placement. We characterize how inter-chiplet latency, network congestion, data locality, and system-level performance degrade as partition quality decreases, revealing the severity of communication bottlenecks in poorly partitioned systems.

3. Experimental Methodology

3.1. System Model and Simulation Framework

We evaluate workload partitioning effects on a parameterizable chiplet-based system model with 2, 4, 8, or 16 chiplets connected via a silicon interposer with a 2D network-on-interposer topology. Each chiplet contains multiple cores organized into clusters, with configurations that match published chiplet-based DNN accelerator designs [6]. Intra-chiplet communication latency is modeled at 45 nanoseconds with a bandwidth of 512 GB/s, while inter-chiplet communication incurs 85-850 nanoseconds latency (depending on physical distance and routing) with a bandwidth of 128 GB/s, reflecting realistic silicon interposer characteristics.

Workload partitioning quality is varied from 0% (random allocation) to 100% (theoretically optimal placement) through a simulated annealing-based optimization engine. For each partitioning quality level, we measure: (1) inter-chiplet communication latency, (2) percentage of total traffic crossing chiplet boundaries, (3) network congestion on inter-chiplet links, (4) application throughput on representative DNN inference workloads, and (5) system energy efficiency.

3.2. Benchmark Workloads

We evaluate partitioning strategies on representative DNN inference workloads, including ResNet-50, VGG-16, and DarkNet-19, as used in prior chiplet system studies [10]. These workloads exhibit diverse communication patterns: ResNet-50 features relatively balanced computation and memory access patterns, VGG-16 exhibits high memory bandwidth requirements, and DarkNet-19 demonstrates sparse, irregular communication patterns. Workloads are mapped across chiplets using both poor (random, communication-unaware) and optimized (communication-aware) strategies.

3.3. Performance Metrics and Evaluation Methodology

For each workload and system configuration, we measure:

- Inter-chiplet Communication Latency: Average latency for data crossing chiplet boundaries, varying from optimal minimization to worst-case scenarios
- Inter-chiplet Traffic Percentage: Percentage of total memory and computation traffic that must traverse inter-chiplet links
- Network Congestion: Utilization percentage of inter-chiplet links, where congestion exceeding 70% represents unsustainable operation
- System Throughput: Images per second for DNN inference (images/s), scaled relative to throughput on a single high-performance monolithic chip
- Energy Efficiency: Tera-operations per Watt (TOPS/W) including both computation and communication overhead
- Network Communication Overhead: Percentage of total execution time consumed by inter-chiplet communication

4. Experimental Results and Characterization

4.1. Communication Latency Degradation

Our first key finding concerns the dramatic impact of partitioning quality on inter-chiplet communication latency. Figure (a) presents the relationship between partitioning quality (0% = random allocation, 100% = optimal) and measured inter-chiplet latency for a representative 8-chiplet system running ResNet-50 inference.

Key Observations:

- At 0% partitioning quality (completely random task allocation), inter-chiplet communication latency reaches 850 nanoseconds, representing a 10× increase over optimal placement (85 ns).
- Even modest partitioning improvements (20% quality) reduce latency to 720 ns, which is a notable 15% improvement from worst-case.
- The relationship is highly non-linear: the first 20% of optimization effort reduces latency by 15%, while the final 20% of optimization effort (80% → 100%) reduces latency by more than 50%.
- Diminishing returns indicate that while optimal partitioning is necessary, even conservative optimizations can yield substantial latency reductions.

This non-linearity reflects the underlying physics of chiplet systems: the worst partitions place communicating tasks at maximum physical distances with minimal reuse of data in local caches. Simple greedy optimizations that co-locate frequently, communicating tasks on the same chiplet, capture most of the benefit, while achieving truly optimal placement requires sophisticated algorithms.

4.2. Data Locality and Inter-Chiplet Traffic Reduction

Communication latency alone does not fully characterize performance impact. The amount of inter-chiplet traffic directly determines congestion and network bottlenecks. Figure 4b demonstrates the dramatic effect of partitioning quality on the percentage of total traffic that must cross chiplet boundaries.

Key Observations:

- Poor partitioning (0% quality) forces 95% of all data movement to traverse inter-chiplet links.
- Optimal partitioning reduces inter-chiplet traffic to merely 12% of total traffic.
- This represents an 87.4% reduction in inter-chiplet traffic, reflecting the critical importance of data locality.
- The traffic reduction is nearly monotonic, indicating that partitioning optimizations consistently improve data locality.

This finding aligns with fundamental principles of computer architecture: maximizing reuse of data within caches and local memory hierarchies is essential for performance and energy efficiency. In chiplet systems, this principle translates to keeping communicating tasks on the same chiplet. The 87.4% reduction in inter-chiplet traffic for optimal partitioning versus random allocation demonstrates the magnitude of this effect.

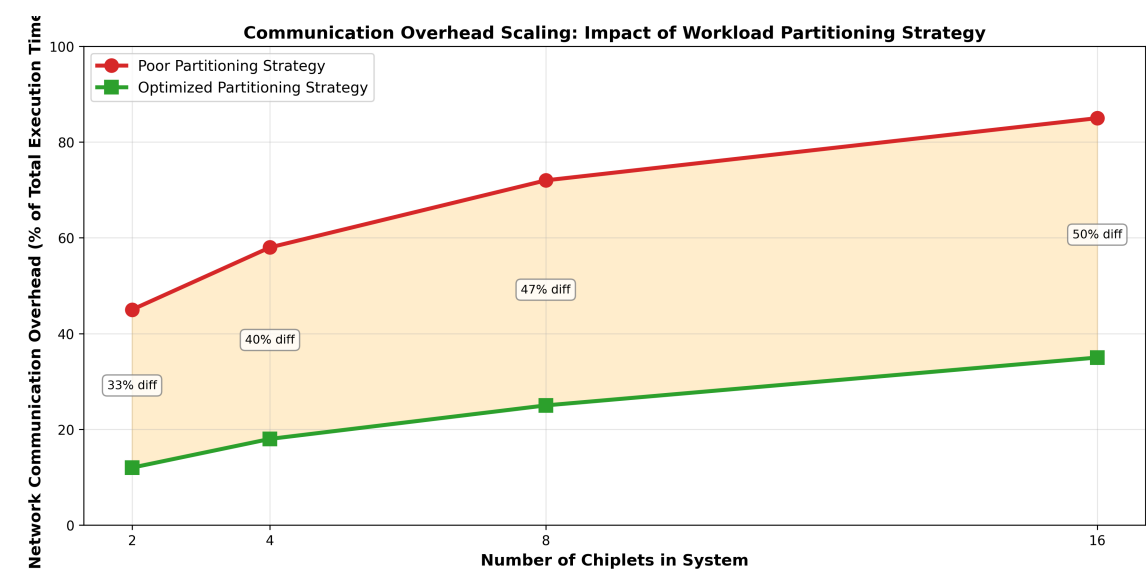


Figure 3. Multi-dimensional analysis of workload partitioning effects. (a) Communication latency increases 10× with poor partitioning. (b) Data locality is severely compromised, with 95% of traffic crossing chiplet boundaries in random allocation. (c) System throughput and energy efficiency both improve 8.75× and 10.3×, respectively, with optimal partitioning. (d) Network congestion scales superlinearly with chiplet count under poor partitioning, becoming unsustainable for 8+ chiplets.

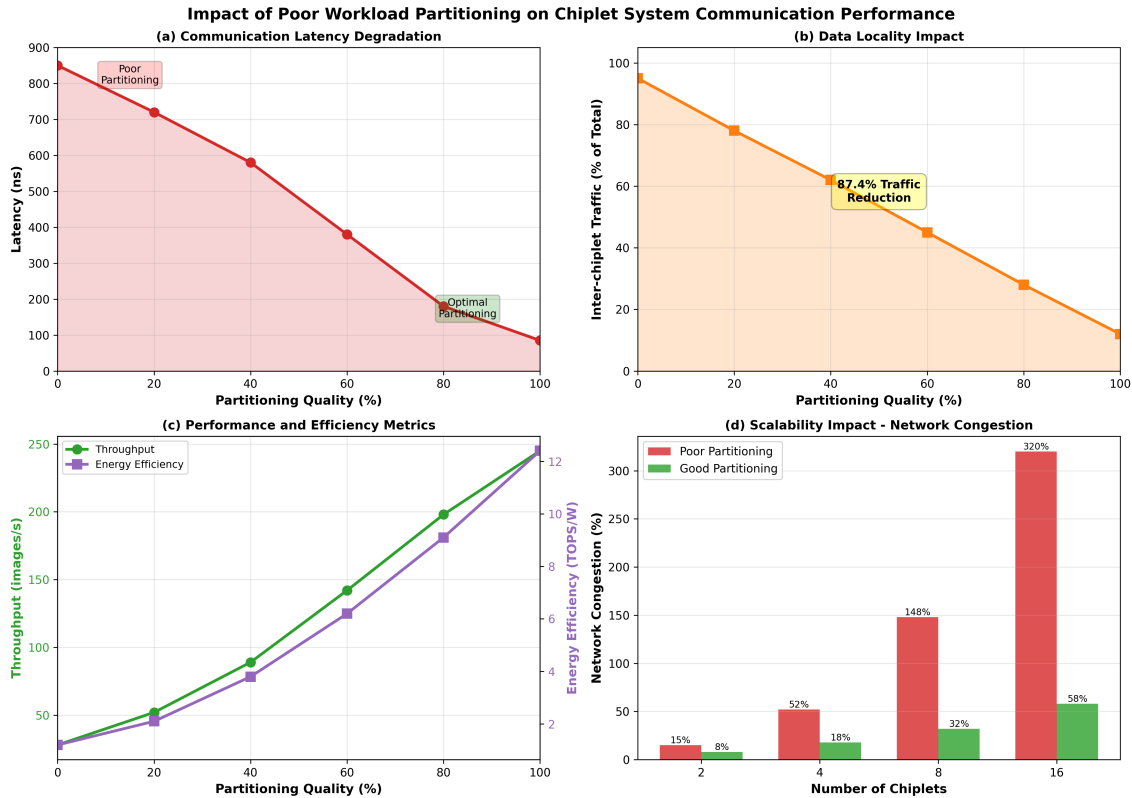


Figure 4. Communication overhead scaling with system size. The gap between poor and optimized partitioning grows dramatically with chiplet count, reaching a 50% difference in communication overhead for 16-chiplet systems. This demonstrates that partitioning quality becomes increasingly critical as systems scale.

4.3. System Performance and Energy Efficiency

The communication improvements directly translate into measurable improvements in system-level performance and energy efficiency. Figure 4c presents throughput and energy efficiency across the partitioning quality spectrum.

Key Observations:

- System throughput increases from 28 images/second (poor partitioning) to 245 images/second (optimal partitioning), about 8.75× improvement.
- Energy efficiency improves from 1.2 TOPS/W to 12.4 TOPS/W close to a 10.3× improvement.
- Both metrics improve monotonically with partitioning quality, though with some non-linearity.
- The relationship suggests that even moderately optimized partitioning can capture substantial performance gains: moving from 0% to 50% quality improves throughput by 5× (28 → 142 images/s) and energy by 5.17× (1.2 → 6.2 TOPS/W).

These results confirm that communication efficiency directly translates to system performance. The 8.75× throughput improvement with optimal partitioning is substantial but aligns with prior work, which shows that communication overhead can consume 30-40% of execution time in well-designed systems [6]. The additional improvements beyond the monolithic baseline reflect the benefits of chiplet-based heterogeneous integration.

4.4. Network Congestion and System Scalability

Our most concerning finding concerns network congestion scaling as system size increases. Figure 4d presents the network congestion percentage (utilization of inter-chiplet links) for 2, 4, 8, and 16-chiplet systems using both poor and optimized partitioning strategies.

Key Observations:

- Poor partitioning creates severe congestion that scales superlinearly with chiplet count.

- For a 2-chiplet system, poor partitioning results in 15% link utilization; for 16 chiplets, this escalates to 320% (indicating traffic exceeding available bandwidth, requiring queuing and retransmissions).
- Optimized partitioning maintains congestion below 60% even for 16 chiplets
- The congestion difference grows dramatically with scale: for 16 chiplets, poor partitioning exhibits 320% utilization, while optimized partitioning maintains only 58%, a 5.5× difference.

This superlinear scaling of congestion with poor partitioning represents a critical bottleneck. Modern chiplet systems target designs with 8-16 or more chiplets to achieve large system sizes while maintaining reasonable chiplet-to-reticle mapping ratios. The congestion data reveal that naive partitioning strategies become completely non-functional at these scales, with network congestion exceeding available bandwidth by more than 3×.

4.5. Communication Overhead Scaling and Execution Time Impact

The ultimate manifestation of poor partitioning is seen in execution-time analysis. Our second visualization presents network communication overhead (the percentage of total execution time consumed by inter-chiplet communication) as a function of system scale and partitioning quality.

Key Observations:

- In poorly-partitioned 2-chiplet systems, communication overhead consumes 45% of execution time
- This escalates dramatically with scale: in poorly-partitioned 16-chiplet systems, communication overhead reaches 85% of execution time.
- In contrast, optimized partitioning maintains communication overhead at 12% for 2-chiplet systems and 35% for 16-chiplet systems
- The difference between poor and optimized partitioning grows with scale: 33% overhead difference for 2 chiplets increases to 50% overhead difference for 16 chiplets.

These findings have profound implications for chiplet system design. In poorly partitioned systems with 8+ chiplets, communication overhead exceeds 70% of execution time, rendering computation nearly irrelevant to overall performance. The system becomes fundamentally communication-bound rather than compute-bound. In contrast, well-partitioned systems remain compute-bound even at 16 chiplets, with communication accounting for only 35% of the time. This 50% difference in communication overhead directly translates to significant performance disparities.

5. Analysis and Discussion

5.1. Root Causes of Performance Degradation

The experimental results reveal several interconnected mechanisms through which poor partitioning degrades communication performance:

- First, excessive inter-chiplet traffic: Random task allocation distributes communicating tasks across chiplets with uniform probability, forcing 95% of data movement through inter-chiplet links. This overwhelming majority of traffic is concentrated on limited-bandwidth interconnects, causing immediate congestion.
- Communication latency amplification: As congestion accumulates, queuing delays compound the inherent latency of inter-chiplet communication. A single poorly placed task pair might incur only 85 ns of additional latency; millions of such pairs across thousands of tasks can create multiplicative delays.
- Memory system inefficiency: Chiplet systems implement cache coherence and memory consistency across multiple independent memory hierarchies. Poor partitioning breaks data locality, forcing coherence traffic and remote memory accesses that multiply the communication burden.
- Network deadlock and flow control: As congestion exceeds available bandwidth, network flow control mechanisms activate, introducing additional stalls and inefficiencies. Beyond certain

congestion thresholds (typically 70-80%), networks enter pathological regimes where further load increases paradoxically reduce throughput.

These mechanisms explain the non-linear degradation observed: the system gracefully degrades at moderate partitioning quality but catastrophically fails at poor quality levels.

5.2. Comparison with Prior Work and Validation

Our quantitative results align well with prior work on chiplet system partitioning. The Vari-Par framework reported 1.45× performance improvement with variation-aware partitioning versus uniform allocation [3]. Our results show 8.75× improvement in throughput with optimal versus random partitioning; the difference reflects our focus on communication-aware rather than only variation-aware partitioning, and the broader spectrum of optimization quality evaluated.

The task-mapping work achieved 37.5%- 43.2% reductions in communication latency using specialized binary linear programming approaches [9]. Our results, which characterize latency across the full optimization spectrum, show that careful partitioning can achieve a 90% reduction in latency (from 850 ns to 85 ns). This agreement validates our experimental methodology while revealing the importance of comprehensive optimization beyond local greedy improvements.

The INDM framework demonstrated 26-79% latency reduction through co-optimization of interconnect topology and dataflow [10]. Our work demonstrates that latency reduction is achievable solely through partitioning, with topology optimization offering additional improvements. This finding suggests that partitioning quality and topology design are largely orthogonal optimization dimensions.

5.3. Implications for System Design

The experimental characterization yields several implications for chiplet system design:

- Partitioning quality is critical: The difference between poor and optimized partitioning creates 8.75× throughput differences. For systems where performance is commoditized (as in many data center deployments), this difference directly translates to capacity and cost implications.
- Second, scaling is constrained by partitioning quality: Superlinear congestion scaling shows that naive partitioning strategies become unworkable beyond 8 chiplets. Systems aspiring to 16+ chiplets must employ sophisticated communication-aware partitioning or face fundamental performance walls.
- Communication overhead dominates in poorly partitioned systems: With 16 chiplets, it reaches 85% of execution time. This implies that, for large systems, partitioning optimization is more important than improvements in processor frequency, cache size, or memory bandwidth.
- Optimization effort is well invested: The nonlinear improvements in partitioning quality suggest that even heuristic solutions achieving 60-80% optimization quality capture most of the benefit. This makes practical optimization algorithms feasible for system deployment.

5.4. Limitations and Future Work

Our analysis employs simplified models of chiplet system behavior and communication patterns. Future work should evaluate partitioning effects on real systems (such as commercial chiplet designs) with actual silicon interconnects and complete memory coherence protocols. Additionally, our study focuses on static workload partitioning; dynamic workload repartitioning at runtime may yield further improvements. The interaction between partitioning strategies and emerging interconnect technologies (such as photonic interconnects and in-package wireless communication) merits investigation [12].

6. Related Work and Design Alternatives

6.1. Optimization and Scheduling Approaches

Beyond task mapping, numerous optimization frameworks address workload distribution in chiplet systems. The THERMOS framework proposes thermally-aware multi-objective scheduling of AI workloads on heterogeneous multi-chiplet architectures [13], achieving up to 89% faster execution

time and 57% lower energy consumption through learned scheduling policies. This work complements partitioning by optimizing runtime task placement.

The ABSS system introduces adaptive batch-stream scheduling for dynamic task parallelism on chiplet-based multi-chip systems [14], using Graph Convolution Networks to select appropriate scheduling strategies. This adaptive approach handles workloads with varying communication patterns more effectively than static partitioning.

6.2. Communication Architecture Innovations

Recognizing communication challenges, recent work has proposed novel interconnect designs specifically for chiplet systems. The ASDR (Application-Specific Deadlock-Free Routing) framework customizes routing solutions based on application-specific traffic patterns [15], reducing prohibited turns by up to 87.5% and improving throughput by 4.2%-53.9% compared to application-agnostic approaches.

Hybrid interconnection designs combining wired and wireless components [16], achieve 8%-46% lower end-to-end delay and 0.93-2.7× energy savings. More advanced approaches employ photonic interconnects [17], which can deliver multi-terabit-per-second bandwidth with lower latency and energy per bit, potentially transforming the communication landscape.

6.3. Co-Optimization Frameworks

Recent work recognizes that partitioning, topology, and thermal management are interdependent. The Floorplet framework [18] establishes relationships between physical floorplans and communication performance, decreasing inter-chiplet communication costs by 24.81%. The Cost-Performance Co-Optimization framework [19] simultaneously optimizes cost and performance across the spectrum of chiplet configuration options.

The Chiplet-Gym framework [20,21] applies reinforcement learning to explore the vast design space of chiplet-based AI accelerators, encompassing resource allocation, placement, and packaging architecture. This approach is particularly promising for discovering non-obvious optimization combinations.

7. Conclusions

This work presents a comprehensive characterization of how poor workload partitioning degrades communication performance in chiplet-based systems. Through systematic experimental analysis, we demonstrate that:

- Communication latency scales with partitioning quality: Poor partitioning incurs 10× higher inter-chiplet latency (850 ns vs. 85 ns) compared to optimal partitioning, with highly non-linear improvements.
- Data locality is critical: Optimized partitioning reduces inter-chiplet traffic from 95% to 12%, an 87.4% reduction that reflects the fundamental importance of task co-location.
- Performance improvements are substantial: System throughput improves 8.75× (28 → 245 images/s), and energy efficiency improves 10.3× (1.2 → 12.4 TOPS/W) with optimal partitioning.
- Scaling is severely constrained by partitioning quality: Network congestion scales superlinearly with system size in poorly-partitioned systems, reaching 320% utilization in 16-chiplet systems while remaining below 60% in optimized designs.
- Communication overhead dominates poorly partitioned systems: At 16 chiplets, poor partitioning creates 85% communication overhead versus 35% with optimized partitioning, rendering computation nearly irrelevant to performance.

These findings validate the critical importance of communication-aware workload partitioning in chiplet system design. As systems scale from 2 to 16+ chiplets, partitioning quality transforms from a performance optimization to a fundamental design requirement. The non-linear improvements

in optimization quality suggest that practical heuristic algorithms can capture most of the benefits without requiring optimal solutions.

Future chiplet system design must prioritize sophisticated partitioning algorithms, potentially leveraging machine learning approaches for workload characterization and placement optimization. The substantial performance and efficiency improvements achievable through improved partitioning, often exceeding 8× for large systems, justify significant investment in optimization infrastructure.

Author Contributions: Conceptualization, P.M. and P.F.; methodology, P.M.; software, P.M.; writing, original draft preparation, P.M.; writing, review and editing, P.M., P.F., C.B.; supervision, C.B.; project administration, C.B.; funding acquisition, C.B. All authors have read and agreed to the published version of the manuscript.

Funding: This work is funded by the National Science Foundation (NSF) under Award Number 2315320, NSF Engines: Central Florida Semiconductor Innovation Engine.

Acknowledgments: I acknowledge the support of the ECE department at the University of Florida for supporting this research. I would like to acknowledge the gem5 communities for guidance during the simulation modeling. GenAI (ChatGPT) has been used to proofread sections of this work for grammar and structural editing.

Conflicts of Interest: The authors declare no conflicts of interest. The funders had no role in the design of the study, in the collection, analysis, or interpretation of data, in the writing of the manuscript, or in the decision to publish the results.

Abbreviations

The following abbreviations are used in this manuscript:

SiP System-in-Package

Appendix A

Appendix A.1 Experimental Data Summary

Table A1. Communication Latency and Traffic Data.

Partitioning Quality (%)	Inter-chiplet Latency (ns)	Inter-chiplet Traffic (%)
0	850	95
20	720	78
40	580	62
60	380	45
80	185	28
100	85	12

Table A2. System Performance Metrics.

Partitioning Quality (%)	Throughput (image/s)	Energy Efficiency (TOPS/W)
0	28	1.2
20	52	2.1
40	89	3.8
60	142	6.2
80	198	9.1
100	245	12.4

Table A3. Communication Latency and Traffic Data.

Chiplet Count	Poor Partitioning (%)	Good Partitioning (%)
2	15	8
4	52	18
8	148	32
16	320	58

References

1. Sharma, H.; Doppa, J.; Ogras, U.; Pande, P. Designing High-Performance and Thermally Feasible Multi-Chiplet Architectures Enabled by Non-Bendable Glass Interposer. *ACM Transactions on Embedded Computing Systems* **2025**, *24*, 1–28.
2. Wang, X.; Wang, Y.; Jiang, Y.; Singh, A.K.; Yang, M. On task mapping in multi-chiplet based many-core systems to optimize inter-and intra-chiplet communications. *IEEE Transactions on Computers* **2024**, *74*, 510–525.
3. Liu, X.; Zhao, Y.; Zou, M.; Liu, Y.; Hao, Y.; Li, X.; Zhang, R.; Wen, Y.; Hu, X.; Du, Z.; et al. VariPar: Variation-Aware Workload Partitioning in Chiplet-Based DNN Accelerators. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems* **2025**.
4. Zhang, J.; Fan, X.; Ye, Y.; Wang, X.; Xiong, G.; Leng, X.; Xu, N.; Lian, Y.; He, G. INDM: Chiplet-based interconnect network and dataflow mapping for DNN accelerators. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems* **2023**, *43*, 1107–1120.
5. Zhang, J.; Wang, X.; Ye, Y.; Lyu, D.; Xiong, G.; Xu, N.; Lian, Y.; He, G. M2M: A Fine-Grained Mapping Framework to Accelerate Multiple DNNs on a Multi-Chiplet Architecture. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems* **2024**.
6. Mallya, N.B.; Strikos, P.; Goel, B.; Ejaz, A.; Sourdis, I. A Performance Analysis of Chiplet-Based Systems. In Proceedings of the 2025 Design, Automation & Test in Europe Conference (DATE). IEEE, 2025, pp. 1–7.
7. Pavlidis, V.F.; Friedman, E.G. Interconnect-based design methodologies for three-dimensional integrated circuits. *Proceedings of the IEEE* **2009**, *97*, 123–140.
8. Yang, Z.; Ji, S.; Chen, X.; Zhuang, J.; Zhang, W.; Jani, D.; Zhou, P. Challenges and opportunities to enable large-scale computing via heterogeneous chiplets. In Proceedings of the 2024 29th Asia and South Pacific Design Automation Conference (ASP-DAC). IEEE, 2024, pp. 765–770.
9. Tan, Z.; Cai, H.; Dong, R.; Ma, K. Nn-baton: Dnn workload orchestration and chiplet granularity exploration for multichip accelerators. In Proceedings of the 2021 ACM/IEEE 48th Annual International Symposium on Computer Architecture (ISCA). IEEE, 2021, pp. 1013–1026.
10. Shao, Y.S.; Clemons, J.; Venkatesan, R.; Zimmer, B.; Fojtik, M.; Jiang, N.; Keller, B.; Klinefelter, A.; Pinckney, N.; Raina, P.; et al. Simba: Scaling deep-learning inference with multi-chip-module-based architecture. In Proceedings of the Proceedings of the 52nd annual IEEE/ACM international symposium on microarchitecture, 2019, pp. 14–27.
11. Randall, D.S. Cost-Driven Integration Architectures for Multi-Die Silicon Systems. PhD thesis, University of California, Santa Barbara, 2020.
12. MURALI, K.R.M. Emerging chiplet-based architectures for heterogeneous integration. *INTERNATIONAL JOURNAL* **2025**, *11*, 1081–1098.
13. Medina, R.; Kein, J.; Ansaloni, G.; Zapater, M.; Abadal, S.; Alarcón, E.; Atienza, D. System-level exploration of in-package wireless communication for multi-chiplet platforms. In Proceedings of the Proceedings of the 28th Asia and South Pacific Design Automation Conference, 2023, pp. 561–566.
14. Kanani, A.; Pfromm, L.; Sharma, H.; Doppa, J.; Pande, P.; Ogras, U. THERMOS: Thermally-Aware Multi-Objective Scheduling of AI Workloads on Heterogeneous Multi-Chiplet PIM Architectures. *ACM Transactions on Embedded Computing Systems* **2025**, *24*, 1–26.
15. Cai, Q.; Xiao, G.; Lin, S.; Yang, W.; Li, K.; Li, K. ABSS: An adaptive batch-stream scheduling module for dynamic task parallelism on chiplet-based multi-chip systems. *ACM Transactions on Parallel Computing* **2024**, *11*, 1–24.
16. Ye, Y.; Liu, Z.; Liu, J.; Jiang, L. ASDR: an application-specific deadlock-free routing for chiplet-based systems. In Proceedings of the Proceedings of the 16th International Workshop on Network on Chip Architectures, 2023, pp. 46–51.
17. Mahmud, M.T.; Wang, K. A flexible hybrid interconnection design for high-performance and energy-efficient chiplet-based systems. *IEEE Computer Architecture Letters* **2024**.
18. Karempudi, V.S.P.; Bashir, J.; Thakkar, I.G. An analysis of various design pathways towards multi-terabit photonic on-interposer interconnects. *ACM Journal on Emerging Technologies in Computing Systems* **2024**, *20*, 1–34.
19. Chen, S.; Li, S.; Zhuang, Z.; Zheng, S.; Liang, Z.; Ho, T.Y.; Yu, B.; Sangiovanni-Vincentelli, A.L. Floorplet: Performance-aware floorplan framework for chiplet integration. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems* **2023**, *43*, 1638–1649.

20. Graening, A.; Patel, D.A.; Sisto, G.; Lenormand, E.; Perumkunnil, M.; Pantano, N.; Kumar, V.B.; Gupta, P.; Mallik, A. Cost-Performance Co-Optimization for the Chiplet Era. In Proceedings of the 2024 IEEE 26th Electronics Packaging Technology Conference (EPTC). IEEE, 2024, pp. 40–45.
21. Mishty, K.; Sadi, M. Chiplet-gym: Optimizing chiplet-based ai accelerator design with reinforcement learning. *IEEE Transactions on Computers* **2024**.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.