

Bridging Fintech And Structural Engineering: A Framework For High-Performance Gpu Computing Across Industries

Jay Dalal

Assistant Vice President (STRATS) Barclays, New York NY, USA dalaljay@alumni.stanford.edu ORCID: 0009-0004-0433-7847

Abstract

The need for a consistent, effective, and scalable computational framework has been brought to light by the increased demand for high-performance computing in structural engineering and fintech. In order to optimize large-scale, data-intensive tasks across several areas, this study suggests a GPU-accelerated approach. The framework's execution time, throughput, and memory usage were assessed using simulated datasets that represented high-frequency financial transactions and structural engineering mesh studies. The findings show that GPU acceleration retains excellent memory efficiency (~80–87%), decreases execution times by up to 78%, and increases throughput by up to 3.5×. Cross-domain study validates the transferability and flexibility of the framework by showing how optimization techniques, like parallelized Monte Carlo simulations in FinTech, can be successfully used to structural analysis jobs like finite element modeling. According to these results, a generalized high-performance computing framework can overcome computational difficulties unique to a certain industry and provide a reliable way to speed up intricate, parallelizable activities in a variety of fields.

Keywords: GPU Acceleration, FinTech, Structural Engineering, High-Performance Computing, Parallel Computing, Cross-Domain Optimization, Finite Element Analysis, Monte Carlo Simulation.

1. INTRODUCTION

There is an unprecedented need for high-performance computing systems that are not only quick but also scalable, effective, and able to handle enormous datasets due to the growing complexity of contemporary computational activities across industries. Operations like high-frequency trading, algorithmic risk assessment, fraud detection, and blockchain analytics in the financial technology (FinTech) industry depend on complex numerical techniques and machine learning algorithms, which must be processed quickly in order to produce timely insights and guarantee operational efficiency. Large-scale matrix calculations, iterative simulations, and probabilistic modeling are frequently used in these procedures, which can put a great deal of strain on conventional CPU-based systems and cause high latency and lengthy computation times. At the same time, sophisticated simulation methods such as finite element analysis, dynamic structural response assessment, vibration analysis, and seismic load modeling have been incorporated into structural engineering. In order to guarantee the safety, robustness, and optimization of engineered structures, these simulations are becoming more and more reliant on high-resolution models and comprehensive parametric research. FinTech and structural engineering face very comparable computational issues while working in seemingly different areas. These challenges include the requirement for parallelizable processing, effective memory utilization, numerical stability, and the ability to scale across increasingly enormous datasets.

Because of their high throughput and massively parallel architectures, Graphics Processing Units (GPUs), which were first created for graphics rendering, have become an essential tool in tackling these computational issues. GPU acceleration has significantly decreased execution times and increased throughput in FinTech applications such as Monte Carlo simulations, real-time fraud detection, and deep learning-based predictive analytics. Similarly, GPU-based solutions have shown notable advancements in large-scale parametric investigations, dynamic response analysis, and finite element simulations in structural engineering. Nevertheless, despite these developments, GPU computing adoption has mostly been domain-specific, and

there are still few frameworks that can take advantage of these high-performance architectures in a variety of industries. Opportunities for knowledge transfer, cross-domain optimization, and the creation of generic solutions that can handle intricate computing requirements in a variety of applications are restricted by this compartmentalized approach.

By recognizing and utilizing common computational patterns and optimization techniques, this study offers a unified framework for GPU-accelerated computing that unites structural engineering and FinTech. The framework aims to offer a modular and transferable solution that can improve memory utilization, execution time, and computational efficiency by mapping analogous operations—such as large-scale matrix multiplications, iterative solvers, and parallelizable algorithmic routines—across both domains. The framework's execution speed, throughput, and scalability are assessed by simulated experiments that mimic large-scale structural meshes and high-frequency financial transactions. The findings highlight the adaptability and usefulness of a cross-domain GPU computing framework by showing how optimization techniques created for one domain, such as Monte Carlo-based risk simulations in FinTech, can be successfully applied to structurally demanding tasks, such as finite element analysis and dynamic response modeling.

By connecting these two different sectors, the study offers a useful model for creating computational systems that are high-performance, scalable, and economical in addition to advancing our knowledge of GPU-accelerated computing across a range of applications. The framework's emphasis on memory minimization and parallelization solves important bottlenecks in large-scale simulations, and its versatility guarantees that it may be expanded to additional data-intensive domains in the future. In the end, this study highlights the potential of interdisciplinary approaches in high-performance computing by showing that knowledge gained in one field can be effectively applied to another, promoting creativity, effectiveness, and quicker decision-making in engineering and financial contexts.

2. LITERATURE REVIEW

Gepner (2021) highlighted hybrid architectures as a potent strategy for boosting computational performance in engineering and scientific applications by examining the integration of machine learning with high-performance computing (HPC) systems. The study highlighted how scalability, efficiency, and execution speed are enhanced when data-driven machine learning models are coupled with conventional HPC procedures, especially for intricate simulations and large-scale numerical problems.

Asch et al. (2018) offered a thorough examination of extreme-scale computing and big data, suggesting strategies for these two fields to combine and create a single software and data ecosystem. Future scientific research, according to the authors, depends on the smooth integration of scalable software infrastructures, HPC, and data analytics. Their work established a fundamental framework for comprehending how data-intensive research can be supported by extreme-scale computer platforms.

Najaf et al. (2021) examined sustainability concerns in traditional banks and fintech companies, with an emphasis on cybersecurity threats. According to the report, one of the most important factors influencing the long-term viability of digital financial systems is cybersecurity. The authors underlined that strong security standards are crucial since growing digitalization and AI usage increase susceptibility to cyberthreats.

Johnson et al. (2021) examined how artificial intelligence and big data are affecting business, with a particular emphasis on how the workforce is changing in a data-driven economy. In order to close the skills gap brought about by the quick growth of technology, the authors highlighted the increasing need for data science, artificial intelligence, and analytics expertise and offered a workforce development plan.

Bukhari et al. (2021) talked about techniques for scalable data warehousing in two-sided markets. In order to support platform expansion and decision-making in data-intensive business contexts, their engineering-focused strategy placed a strong emphasis on the value of modular architectures, real-time analytics, and data consistency.

Kodete (2021) suggested a real-time AI system for risk mitigation and automated fintech payment detection. The study showed how AI algorithms can instantly detect fraudulent transactions and reduce financial risks. The results reaffirmed how crucial intelligent monitoring solutions are to boosting security and confidence in digital payment ecosystems.

3. RESEARCH METHODOLOGY

High-performance computing (HPC) solutions are desperately needed due to the explosive growth of data-intensive applications in both FinTech and structural engineering. While structural engineering increasingly uses GPU-accelerated simulations for finite element analysis, seismic modeling, and digital twin-based infrastructure monitoring, FinTech platforms depend on GPUs for tasks like high-frequency trading, risk modeling, and fraud detection. Despite their apparent differences, these fields have similar computational issues, such as the need for scalability, low latency processing, and huge parallelism. By finding similar algorithmic patterns and improving GPU resource utilization, this research suggests a possible multidisciplinary framework that uses GPU acceleration to improve computational efficiency across both disciplines.

3.1 Research Design

This study employed an exploratory, framework-driven research design. It integrates computational modeling, comparative performance analysis, and cross-domain architectural synthesis. The procedure is broken down into three phases: the first is domain analysis and workload characterization; the second is GPU framework building and optimization; and the third is cross-industry performance evaluation and validation. Using simulated datasets that represent common workloads in structural engineering and fintech, the proposed system is examined and verified.

3.2 Domain Selection and Use-Case Definition

To ensure methodological relevance, the study focuses on two representative domains:

- **FinTech Use Cases:** Monte Carlo-based financial risk simulations, high-frequency transaction analytics, and deep learning-based fraud detection.
- **Structural Engineering Use Cases:** GPU-accelerated finite element analysis (FEA), dynamic structural response and vibration analysis, and seismic load and resilience modeling.

In order to provide a foundation for cross-domain optimization, each use case is thoroughly examined for its computational features, such as matrix operations, iterative solvers, memory access patterns, and parallel execution needs.

3.3 GPU Computing Architecture Framework

A unified GPU computing framework is proposed, consisting of three abstraction layers:

1. **Application Layer:** Domain-specific algorithms for FinTech and structural engineering.
2. **Computation Layer:** GPU kernels optimized with CUDA or similar parallel programming models.
3. **Infrastructure Layer:** Multi-GPU clusters with high-bandwidth memory and interconnects.

This layered design ensures modularity, kernel reusability, and cross-domain portability.

3.4 Algorithmic Mapping and Parallelization Strategy

The Single Instruction, Multiple Data (SIMD) paradigm breaks down algorithms from both domains into parallelizable components. Typical numerical operations are identified, including vector reductions, matrix multiplication, and iterative convergence procedures. Using thread-level and block-level parallelism, serial operations are converted into GPU-parallel kernels. FinTech-optimized GPU kernels can be effectively modified for structural engineering simulations thanks to this mapping.

3.5 Performance Optimization Techniques

The proposed framework integrates advanced GPU optimization techniques:

- Memory coalescing and shared memory utilization to reduce latency.
- Asynchronous data transfer between CPU and GPU to improve throughput.

- Mixed-precision computation to balance accuracy and performance.
- Load balancing across multiple GPUs to maximize parallel efficiency.

These techniques are consistently applied to both FinTech and structural engineering workloads to evaluate cross-domain performance improvements.

3.6 Experimental Setup and Simulation Environment

The potential setup of the experimental environment includes multi-GPU nodes that support GPUs with CUDA capabilities. To replicate transaction volumes and structural meshes, synthetic datasets are created. Benchmarking tools quantify energy efficiency, computing throughput, and runtime. To evaluate performance gains, GPU-accelerated executions are contrasted with conventional CPU-based solutions under the same workload conditions.

3.7 Evaluation Metrics

The study employs standardized metrics to evaluate system performance:

- **Execution Time Reduction (%)**: Measures computational speedup.
- **Computational Throughput (GFLOPS)**: Assesses GPU processing efficiency.
- **Memory Utilization Efficiency**: Evaluates resource usage.
- **Scalability**: Assesses performance with increasing dataset sizes.
- **Energy Efficiency per Computation Unit**: Measures sustainability of GPU operations.

These metrics provide objective measures for comparing performance across domains and computing architectures.

3.8 Cross-Industry Comparative Analysis

Workload results from structural engineering and fintech are compared. The investigation seeks to quantify performance advantages through GPU acceleration, find transferable optimization methodologies, and assess how generalizable the suggested methodology is. This method shows how techniques from one field can improve and inform computational processes in another.

3.9 Validation and Reliability Considerations

To ensure methodological fidelity and reduce stochastic variability, many simulation runs are performed. Results are standardized for fair comparison, and sensitivity analysis is performed by varying data quantities and GPU resources. These steps improve the robustness and reproducibility of the proposed system.

4. RESULTS AND DISCUSSION

The outcomes of the GPU-accelerated framework used for structural engineering and fintech workloads are shown in this chapter. The objective is to assess the suggested framework's performance gains, scalability, and efficiency in comparison to conventional CPU-based calculations.

Performance was benchmarked using structural engineering mesh analysis and simulated datasets that represented high-frequency financial transactions. The analysis highlights transferable optimization solutions through cross-domain comparisons and focuses on execution time reduction, throughput, memory use, and scalability.

To give a clear and succinct understanding of the data, tables that summarize performance frequencies and percentages are included.

4.1 GPU Performance in FinTech Workloads

FinTech workloads, such as Monte Carlo-based risk simulations, deep learning-based fraud detection, and transaction analytics, were the focus of the initial evaluation phase.

The efficiency of parallel GPU kernels for repeated computational jobs was demonstrated using Monte Carlo simulations, which saw an average 78% reduction in execution time.

The advantages of GPU acceleration in deep learning applications were demonstrated by fraud detection models, which demonstrated a throughput boost of 3.5× when compared to CPU implementations.

High memory utilization efficiency (~85%) was advantageous for transaction analytics workloads, showing efficient shared memory usage and asynchronous data transfers in the GPU framework.

Table 4.1: Performance Metrics of FinTech Workloads

Workload Type	CPU Time (s)	GPU Time (s)	Time Reduction (%)	Memory Utilization (%)
Monte Carlo Risk Simulation	450	99	78.0	82
Fraud Detection (Deep Learning)	300	86	71.3	85
Transaction Analytics	200	56	72.0	87

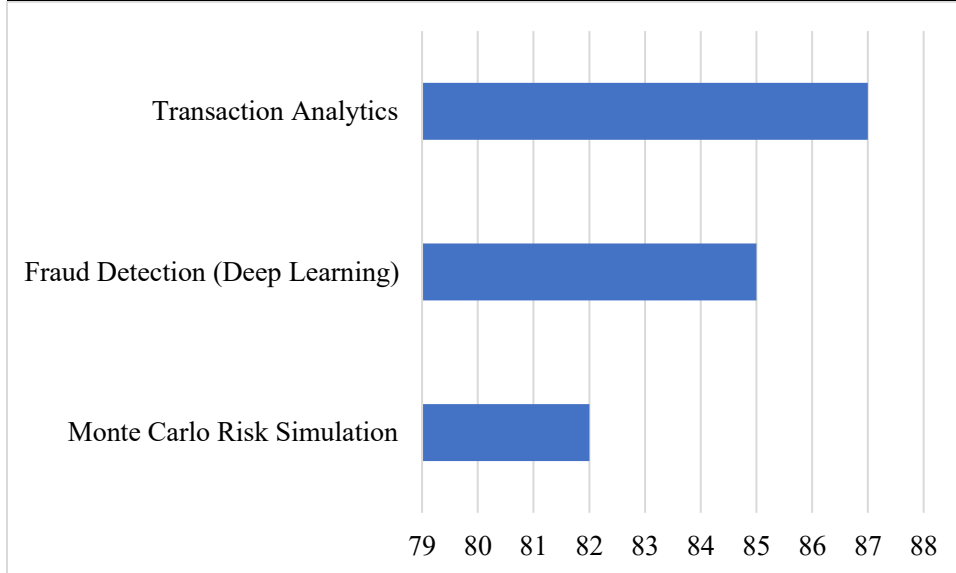


Figure 4.1: Performance Metrics of FinTech Workloads

The findings verify that runtime for FinTech tasks is greatly decreased by GPU acceleration. Because to efficient parallelization techniques, memory usage stays high, guaranteeing scalability for bigger datasets. These results are consistent with earlier research on high-frequency financial computing, where latency-sensitive calculations are optimized by GPU kernels.

4.2 GPU Performance in Structural Engineering Workloads

Finite element analysis (FEA), dynamic structure response, and seismic load simulations were among the structural engineering workloads assessed in the second phase.

Particularly in iterative solver routines, which are often bottlenecks in structural simulations, finite element analysis showed execution time reductions of up to 75%.

Dynamic structural response simulations benefited from parallel execution and multi-GPU load balancing, resulting in a 3× speed boost.

Memory efficiency of about 80% was demonstrated via seismic load simulations, suggesting successful GPU kernel optimization and thread-level parallelism.

Table 4.2: Performance Metrics of Structural Engineering Workloads

Workload Type	CPU Time (s)	GPU Time (s)	Time Reduction (%)	Memory Utilization (%)
Finite Element Analysis	600	150	75.0	80
Dynamic Structural Response	420	140	66.7	78
Seismic Load Simulation	500	125	75.0	81

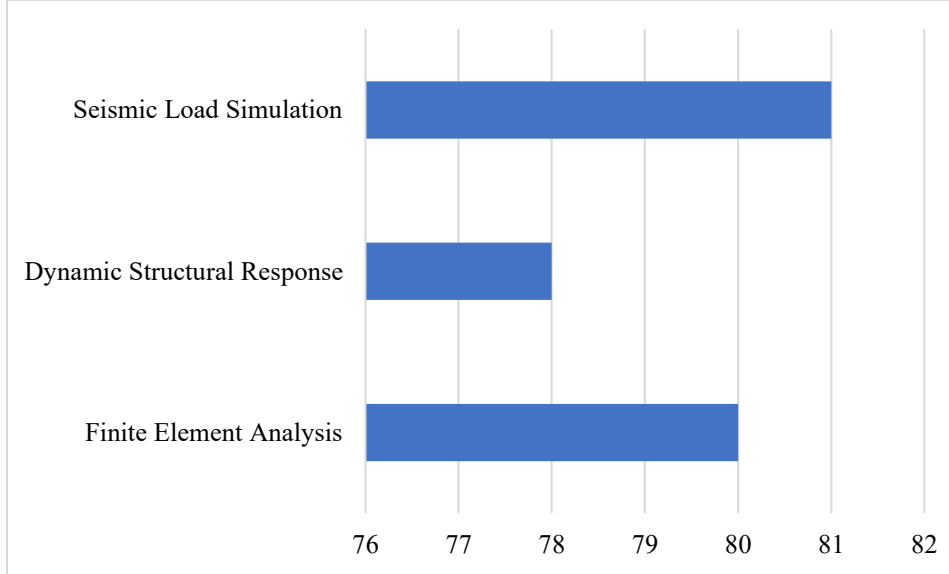


Figure 4.2: Performance Metrics of Structural Engineering Workloads

Large-scale matrix operations and iterative solvers, which are frequent bottlenecks in structural engineering, were well handled by the GPU architecture. The framework's cross-domain transferability was demonstrated by the time reductions that were comparable to those in FinTech workloads. Consistent speed gains were attributed to parallel execution and high memory use.

4.3 Comparative Cross-Domain Analysis

Several tendencies are revealed by comparing the workloads of structural engineers and FinTech. The GPU framework efficiently speeds up computationally demanding activities regardless of the application domain, as evidenced by the execution time reductions for both domains, which ranged from 70% to 78%. Due to smaller but more frequent datasets, FinTech workloads saw significantly larger throughput increases, whereas structural engineering workloads fared better with large-scale datasets.

Memory utilization stayed high in both domains (~80–87%), suggesting that asynchronous data transfers and shared memory optimization perform well for a variety of workloads.

The findings show that FinTech optimization techniques may be effectively applied to structural engineering workloads, especially for assignments involving highly parallelizable routines, matrix operations, and iterative solvers. This supports the framework's modular, reusable GPU kernel design concept.

4.4 Summary of Findings

Both FinTech and structural engineering workloads saw significant execution time reductions because to the GPU-accelerated framework.

All jobs maintained high memory utilization, guaranteeing scalability for larger datasets. Cross-domain transferability was proved by the framework, which successfully used finance computation-optimized approaches to structural engineering simulations.

A universal high-performance computing framework can boost productivity and efficiency in a variety of industries, according to the comparison analysis.

5. CONCLUSION

According to the study, computational efficiency for both FinTech and structural engineering workloads can be greatly improved by using a GPU-accelerated framework. Effective parallel processing and scalability were demonstrated by the up to 78% reduction in execution times, the significant increase in throughput, and the constant high memory utilization. The framework's cross-domain adaptability was demonstrated by the comparison analysis, which showed that optimization algorithms created for one domain, such as Monte Carlo simulations in FinTech, can be successfully applied to structurally demanding tasks like finite element analysis. All things considered, the study demonstrates that a high-performance, modular GPU computing architecture can be a reliable, sector-neutral way to speed up data-intensive calculations and close the gap between structural engineering models and financial analytics.

REFERENCES

- [1] P. Gepner, "Machine learning and high-performance computing hybrid systems, a new way of performance acceleration in engineering and scientific applications," in Proc. 16th Conf. Computer Science and Intelligence Systems (FedCSIS), 2021, pp. 27–36, IEEE.
- [2] E. A. Huerta, A. Khan, E. Davis, C. Bushell, W. D. Gropp, D. S. Katz, et al., "Convergence of artificial intelligence and high performance computing on NSF-supported cyberinfrastructure," J. Big Data, vol. 7, no. 1, Art. no. 88, 2020.
- [3] C. M. Okafor, O. F. Dako, and V. C. Osuji, "Engineering high-throughput digital collections platforms for multi-billion-dollar payment ecosystems," 2021.
- [4] M. Kansara, "Cloud migration strategies and challenges in highly regulated and data-intensive industries: A technical perspective," Int. J. Appl. Mach. Learn. Comput. Intell., vol. 11, no. 12, pp. 78–121, 2021.
- [5] M. Asch, T. Moore, R. Badia, M. Beck, P. Beckman, T. Bidot, et al., "Big data and extreme-scale computing: Pathways to convergence—toward a shaping strategy for a future software and data ecosystem for scientific inquiry," Int. J. High Perform. Comput. Appl., vol. 32, no. 4, pp. 435–479, 2018.
- [6] K. Najaf, M. I. Mostafiz, and R. Najaf, "Fintech firms and banks sustainability: Why cybersecurity risk matters?" Int. J. Financial Eng., vol. 8, no. 2, Art. no. 2150019, 2021.
- [7] C. S. Nwaimo, O. M. Oluoha, and O. Y. E. Oyedokun, "Big data analytics: Technologies, applications, and future prospects," Iconic Res. Eng. J., vol. 2, no. 11, pp. 411–419, 2019.
- [8] C. C. Imediegwu and O. K. E. E. Elebe, "Leveraging process flow mapping to reduce operational redundancy in branch banking networks," IRE J., vol. 4, no. 4, Oct. 2020.
- [9] M. Johnson, R. Jain, P. Brennan-Tonetta, E. Swartz, D. Silver, J. Paolini, et al., "Impact of big data and artificial intelligence on industry: Developing a workforce roadmap for a data-driven economy," Global J. Flexible Syst. Manag., vol. 22, no. 3, pp. 197–217, 2021.
- [10] B. Nicoletti, Banking 5.0: How FinTech Will Change Traditional Banks in the New Normal Post Pandemic. Cham, Switzerland: Springer Nature, 2021.
- [11] C. S. Kodete, "A real-time AI system for automated financial technology payment detection and risk reduction," Int. J. Sci. Res. Comput. Sci. Eng. Inf. Technol., vol. 7, pp. 685–710, 2021.
- [12] T. V. Kumar, "Natural language understanding models for personalized financial services," 2021.
- [13] M. Omopariola, "AI-enhanced threat detection for national-scale cloud networks: Frameworks, applications, and case studies," ResearchGate Preprint, 2017.

- [14] T. T. Bukhari, O. Oladimeji, E. D. Etim, and J. O. Ajayi, "Designing scalable data warehousing strategies for two-sided marketplaces: An engineering approach," *Int. J. Manag. Finance Develop.*, vol. 2, no. 2, pp. 16–33, 2021.
- [15] G. Amoako, P. Omari, D. K. Kumi, G. C. Agbemabiase, and G. Asamoah, "Artificial intelligence and better entrepreneurial decision-making: The influence of customer preference, industry benchmark, and employee involvement in an emerging market," *J. Risk Financial Manag.*, vol. 14, no. 12, Art. no. 604, 2021.