

Parallel Computing Data Processing: Frameworks, Implementations, and Case Studies

Prudhvi Naayini

Independent Researcher

Date of Submission: 01-04-2025

Date of Acceptance: 10-04-2025

ABSTRACT—Parallel computing has become fundamental in processing the massive data volumes generated in modern science and industry. This paper presents a comprehensive survey and practical review of parallel computing in data processing, examining key frameworks (MPI, Open MP, CUDA, Hadoop MapReduce, Apache Spark, etc.), implementations, and real-world case studies. We discuss the architectures underlying shared-memory and distributed-memory systems and illustrate how parallelism on CPUs and GPUs is exploited for high-performance computing (HPC) and big data analytics. We review theoretical foundations (including Amdahl's law for speedup limits) and compare programming models in terms of design, performance, and fault tolerance. A two-column IEEE-style format is used to present critical insights from scholarly literature, with diagrams illustrating parallel architectures and workflows, and tabular summaries of tools and results. Case studies—from scientific HPC simulations to large-scale data analytics—highlight practical parallelism on multi-core CPUs, GPUs, and computing clusters. We conclude with a discussion on the convergence of HPC and big data paradigms, the challenges of heterogeneous computing, and emerging trends. This survey serves as a reference for PhD researchers and practitioners seeking to understand frameworks and real-world practices in parallel computing for data processing.

Index Terms—Parallel computing, Data processing, High-performance computing, Big data analytics, Distributed systems, Multi-core, GPU computing, MPI, Apache Spark

I. INTRODUCTION

High-performance parallel computing is crucial for processing today's enormous data sets and complex computational problems. Future

exascale systems will provide hundreds of times the compute capacity of today's petascale super computers, enabling advanced scientific simulations and big data analytics at unprecedented scale. At the same time, modern "big data" projects are producing massive volumes of information. For example, the Square Kilometre Array (SKA) radio telescope is expected to generate on the order of five zetta bytes of data per year, making it perhaps the largest big data project ever.

Handling such extremes of computation and data requires parallel processing on clusters of many nodes and the efficient use of multi-core CPUs and accelerators (GPUs). Parallel computing involves using multiple processors or computers simultaneously to solve a problem faster. Unlike traditional serial computing where instructions execute one after another, parallel computing breaks a problem into parts that can be executed concurrently.

HPC typically focuses on executing large computational loads (e.g., simulations with heavy floating-point calculations), whereas big data computing centers on handling very large datasets (terabytes to exabytes) with complex data processing. HPC jobs often run on specialized super computers with low-latency interconnects, while big data analytics commonly use distributed clusters of commodity hardware. These two paradigms are beginning to converge, as both scientific computing and data analytics require scalable, fault-tolerant, and efficient parallel processing.

Parallel computing is essential to modern data processing tasks due to the exponential growth of data. By leveraging parallelism, tasks are completed faster and more efficiently, addressing computational bottlenecks in large-scale systems. Parallel computing frameworks such as MPI, OpenMP, CUDA, Hadoop MapReduce, and

Apache Spark provide mechanisms to implement scalable parallelism.

In this paper, we survey the landscape of parallel computing for data processing, covering fundamental concepts, widely- used frameworks, and practical case studies. In Section II (Background), we introduce parallel computing models, memory architectures, and theoretical limits to speed up. In Section III (Frameworks and Tools), we review major programming models and frameworks—ranging from classic HPC libraries (MPI, OpenMP) and GPU programming (CUDA) to big data platforms (Hadoop MapReduce, Spark)—with comparisons of their architectures and use cases. We include illustrative diagrams, such as parallel architecture schematics and data processing workflows, and a summary table of tools. In Section IV (Case Studies), we highlight real-world applications, including scientific HPC simulations and big data analytics tasks, to demonstrate how these frameworks achieve high performance in practice. In Section V (Discussion), we discuss the trade-offs between different approaches, the impact of hardware trends (e.g. GPU acceleration and heterogeneous computing), and the ongoing convergence of HPC and big data ecosystems. Finally, Section VI concludes the paper with insights and future outlook. Throughout, we integrate references to scholarly work to provide a rigorous foundation for this survey.

II. BACKGROUND

A. Parallel Computing Models

Virtually all modern computing systems exploit parallelism at some level. A parallel computer can be a single machine with multiple cores or processing units, or a networked cluster of many machines. Parallelism is inherent in hardware architectures ranging from multi-core desktops to large-scale supercomputers and cloud data centers. Broadly, parallel computing models are classified based on memory architecture and communication mechanisms:

- **Shared-Memory Systems:** Multiple processors access a common physical memory. Threads are typically used to achieve parallelism, with synchronization primitives managing data consistency.
- **Distributed-Memory Systems:** Each processor has its own local memory, and processors communicate through message passing over a network. Clusters and super computers are common examples.
- **Hybrid Systems:** Modern HPC platforms often combine shared-memory multiprocessors (within a node) and distributed memory

(between nodes), requiring hybrid programming approaches (e.g., MPI + OpenMP).

Selecting an appropriate model depends on the problem size, hardware architecture, and communication over head. The choice impacts both performances cal ability and programming complexity.

There are two primary parallel hardware architectures: shared-memory systems, where multiple processors access a common memory, and distributed-memory systems, where each processor or node has its own local memory and communicates with others via a network. Figure 1 illustrates these architectures. In a shared-memory machine, such as a multi- core desktop or an SMP server, all CPUs operate on a single address space (Figure 1, left). In contrast, a distributed cluster consists of nodes each with private memory, requiring message passing for data exchange (Figure 1, right).

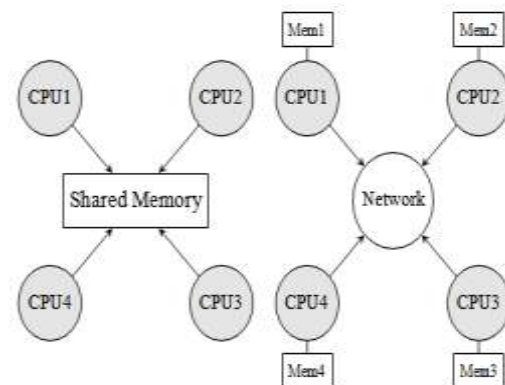


Fig. 1. Illustration of shared vs. distributed memory architectures. **Left:** Multiple processors (CPU1–CPU4) share a single main memory, cooperating via threads. **Right:** Each processor has its own memory (Mem1–Mem4); data exchange occurs over a network interconnect.

More generally, for a given number of processors N , Amdahl's Law is:

$$\frac{1}{(1-f)+\frac{f}{N}} \quad (1)$$

where $0 \leq f \leq 1$. For example, if 90% of a workload can be parallelized ($f=0.9$), the ideal speedup with $N=10$

o-processors is $S(10) \approx 5.3$, and the theoretical maximum as $N \rightarrow \infty$ is $S_{\max}=10$.

This illustrates diminishing returns:

adding processors yields smaller benefits once the serial fraction dominates. Figure 2 conceptually shows how speedup curves level off even as N increases for different values of f . Nonetheless, many problems (especially data-intensive ones) can achieve near-linear speedups by increasing the problem size with the number of processors, an approach described by Gustafson's law (also known as weak scaling).

Parallel computing not only accelerates processing but is also necessary for problems that are too large to fit in the memory or disk of a single machine. Modern big data analytics and scientific simulations both exemplify this: a single computer cannot realistically store or process peta bytes of data or perform billions of computations per second, but a coordinated parallel system can. For instance, the current top supercomputers perform over 10^{18} operations per second (exaflops) on real-world applications.

Such performance is achieved by harnessing tens of thousands of GPUs and CPUs in parallel. On the data side, Internet companies and scientific experiments distribute large datasets across clusters and use parallel data processing frameworks to query and analyze the data efficiently. Another important consideration is fault tolerance and reliability. Large parallel systems with many components have higher chances of hardware or software failures during execution. Traditional HPC jobs typically rely on checkpoint/restart mechanisms (periodically saving state to stable storage) to recover from failures. In contrast, big data processing frameworks often build fault tolerance into the programming model (for example, through replicated data or deterministic recomputation), as will be discussed in later sections. The balance between performance and resilience is a recurring theme in parallel computing. In summary, the background concepts of parallel computing include understanding the hardware architecture (shared vs distributed memory), the theoretical limits to scalability (Amdahl's law and others), and the practical aspects of achieving high performance (efficient communication, workload balance, and fault handling). With this foundation, we next examine the major frameworks and tools that embody these concepts for real-world data processing [1].

III. FRAMEWORKS AND TOOLS

Parallel computing can be implemented via various programming models and frameworks, each suited to different hardware and problem types. This section reviews prominent frameworks spanning HPC and big data domains, highlighting

their architectures, capabilities, and typical use cases.

A. Message Passing Interface (MPI)

MPI is a specification for communication among processes that execute a program on a distributed-memory system. MPI provides routines for sending and receiving messages (point-to-point where $0 \leq f \leq 1$). For example, if 90% of a workload can be parallelized ($f=0.9$), the ideal speedup with $N=10$ to-point communication as well as collective operations (broadcast, reduce, etc.) across processes [2]. MPI programs typically follow the Single Program Multiple Data (SPMD) model, launching identical processes that coordinate through messages. MPI is an HPC library specification for message passing used to program shared and distributed memory systems with bindings for C, C++, and Fortran. Both point-to-point and collective communications are supported.

Over the past decades, MPI has become the de facto standard for parallel programming on clusters. Implementations like MPICH and OpenMPI are available on virtually all HPC platforms. MPI gives programmers fine-grained control over data distribution and communication, enabling high performance on large supercomputers. However, developing MPI code can be complex, since the programmer must manage data locality and handle errors in message exchanges. MPI is widely used in scientific simulations, such as climate modeling and computational fluid dynamics, which run on thousands of processors.

B. OpenMP

OpenMP is an API that supports multi-threaded parallelism in shared-memory environments. It allows developers to parallelize code by inserting compiler directives (pragmas) into C, C++, or Fortran code to specify parallel regions, loops, and work-sharing among threads. The OpenMP API supports multi-platform shared-memory parallel programming in C/C++ and Fortran. The OpenMP API defines a portable, scalable model with a simple and flexible interface for developing parallel applications from the desktop to the supercomputer [3].

In practice, a programmer can mark a loop to run in parallel across threads, and the OpenMP runtime will manage creating threads, dividing the loop iterations, and synchronizing at the end. OpenMP is excellent for incrementally parallelizing computational kernels in an application – one can start with a working serial program and parallelize compute-intensive parts

step by step. Because it uses threads, OpenMP is limited to a single shared-memory node (it does not work across multiple machines without an additional layer like MPI). OpenMP is often combined with MPI for hybrid programming on clusters (using MPI between nodes and OpenMP within each node). Its ease of use makes it a popular choice for leveraging multi-core CPUs. Modern OpenMP also supports directives for offloading computations to accelerators (GPUs), though with less fine control than explicit GPU programming.

C. CUDA and GPU Computing

Graphics Processing Units (GPUs) have evolved into powerful parallel processors used beyond graphics, notably in scientific computing and machine learning. NVIDIA's CUDA is a parallel computing platform and programming model that gives direct access to the GPU's thousands of cores. GPUs follow a Single-Instruction Multiple-Thread (SIMT) model, where thousands of threads execute the same instructions on different data elements in parallel. Rather than processing tasks serially, as a CPU does, a GPU breaks up tasks and runs them in parallel. GPUs have many more cores than CPUs, although they are smaller. With the additional cores, GPUs can handle more calculations at once with greater efficiency, while CPUs, as general-purpose components, are restricted.

For example, a high-end consumer GPU may have over 16,000 CUDACores, compared to 8–32 cores in a typical CPU [4]. CUDA provides C/C++ extensions and libraries to write kernels (functions executed by many threads on the GPU). Memory transfers between CPU and GPU must be managed, and algorithms need to have a high degree of data parallelism to fully utilize the GPU. When they do, speedups of one to two orders of magnitude over CPU-only execution are common for appropriate workloads (e.g., matrix operations, vectorizable loops). Alternatives to CUDA include OpenCL (an open standard for GPU programming) and directive-based approaches like OpenACC. GPUs are now integral to many HPC systems; top supercomputers such as ORNL's Frontier achieve exascale performance with the help of thousands of GPUs and are also used to accelerate big data tasks (for instance, SQL query processing or deep learning training).

D. Map Reduce and Hadoop

In the realm of big data, the MapReduce programming model provided a fundamental paradigm for distributed processing on large clusters. Google's seminal Map Reduce frame-

work introduced a simple model where a dataset is processed by user-defined Map functions that output intermediate key-value pairs, which are then grouped and passed to Reduce functions to produce final outputs. Recent efforts have integrated GPU computing into MapReduce workflows for multi-data sensor fusion and high-performance data handling [5]. Apache Hadoop MapReduce is an open-source implementation of this model. Hadoop Map Reduce is a sub project of Apache Hadoop that provides an implementation of the MapReduce programming model for large-scale data processing [6],[7]. A MapReduce program consists of fine-grained Mapper and Reducer tasks. Both tasks use key-value pairs as input and output. The Mapper transforms input pairs into intermediate pairs, and the Reducer aggregates the intermediate pairs into output. The main features of the MapReduce model are automatic load balancing and recovery from failed tasks (failed tasks are re-executed automatically).

Hadoop MapReduce jobs run on clusters managed by a resource scheduler (YARN in Hadoop). The input data (often stored on the Hadoop Distributed File System, HDFS) is split into blocks across the cluster. As shown in Figure 2, a MapReduce workflow will read these data splits in parallel by launching many map tasks, shuffle and sort the intermediate data by key, and then run reduce tasks to produce output results. The framework handles scheduling tasks on cluster nodes, moving the code to where the data resides (data locality), and dealing with node failures by retrying tasks on other nodes.

This model was highly successful for batch processing of very large datasets (such as aggregating web indexes or log files). However, Map Reduce writes all intermediate results to disk, which can be a bottleneck for iterative algorithms or interactive analytics.

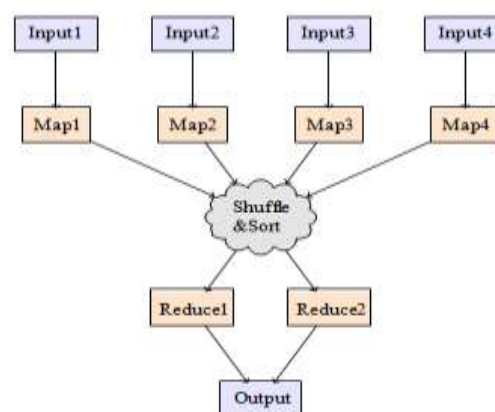


Fig.2.Simplified Map Reduce data processing workflow

As shown in Figure 2 A large input dataset is split into chunks which are processed in parallel by Map tasks. The intermediate results (key–value pairs) from all mappers are shuffled (grouped by key and sorted) before being passed to Reduce tasks that produce the final output. The framework automatically handles task scheduling, data movement (shuffle), load balancing, and fault tolerance (re-executing failed tasks)

E. Apache Spark and Resilient Distributed Datasets(RDDs)

In essence, an RDD is a fault-tolerant, parallel collection of records, which can be operated on in parallel. Spark’s API provides a rich set of operations (e.g., map, filter, join, group-by) that are higher-level than MapReduce, making it easier to compose complex analytics. Crucially, RDDs support in-memory computing; they can persist data in memory across iterations. This is beneficial for algorithms like machine learning iterative optimization or interactive data exploration, where data is reused. Spark only writes

to disk when necessary, reducing the overhead that hampered Hadoop.

Recent frameworks continue to evolve. For example, Merak introduces efficient distributed DNN training with automated 3D parallelism [8], while Graph Tensor optimizes GNN- acceleration for large-scale datasets [9].

Fault tolerance in Spark is achieved differently than in Hadoop. Instead of replicating data to handle failures, Spark tracks the lineage of each RDD—i.e., the sequence of transformations that created it. RDDs provide an interface based on coarse-grained transformations (e.g., map, filter, and join) that apply the same operation to many data items. This allows them to efficiently provide fault tolerance by logging the transformations (lineage) rather than the actual data. If a partition of an RDD is lost, the RDD has enough information about how it was derived from other RDDs to recompute that partition quickly and without costly replication.

TABLE I REPRESENTATIVE PARALLEL COMPUTING FRAMEWORKS FOR DATA PROCESSING

Framework	Parallel Paradigm	Typical Use Case
MPI	Message passing (distributed memory)	HPC simulations on clusters (scientific computing)
OpenMP	Shared-memory threads	Multi-core parallel is min apps (scientific, engineering)
CUDA	Many-core GPU parallelism	GPU-accelerated computing (HPC, AI training)
Hadoop Map Reduce	Batch processing with map/reduce	Big data batch ETL and analytics (disk-based)
Apache Spark	In-memory distributed computing	Big data iterative analytics, machine learning

This design, coupled with lazy evaluation, means Spark can recover from node failures by recomputing missing RDD partitions on another node, while also optimizing the overall workflow execution. In practice, Spark programs are written in languages like Scala, Java, or Python and run on a cluster with a central driver coordinating multiple executor processes (which perform the tasks on data partitions). Spark’s runtime schedules tasks and tries to maintain locality (sending tasks to the nodes that hold the data partitions). Spark also generalizes the MapReduce model: a job may consist of an arbitrary directed acyclic graph (DAG) of operators, not just one map and one reduce phase.

Spark has been shown to greatly outperform Hadoop MapReduce in iterative processing and in-memory analytics— often by 10–100× faster for suitable applications.

It has become a dominant platform for big data analytics, supporting libraries for SQL queries (Spark SQL), machine learning (MLlib), stream processing (Spark Streaming), and graph analytics (GraphX) under a unified engine. One important note is that Spark, by default, is designed for CPU clusters (it uses multi-threading within each executor to exploit multiple cores). However, there have been efforts to integrate GPU acceleration into Spark-like frameworks. For example, HeteroSpark and GPU-enriched Spark variants allow certain operations to run on GPUs,

combining GPU's speed with Spark's scalability.

These efforts manage data partitioning and transfers to ensure GPUs can be used without exposing complexity to the user. Such developments indicate the increasing interest in marrying HPC accelerators with big data frameworks.

Each framework in Table I has its strengths. MPI offers ultimate flexibility and performance on distributed systems but requires manual management of communications. OpenMP is simple for shared-memory multi core machines but cannot scale beyond one node. CUDA unlocks tremendous performance on GPUs for highly parallel tasks, at the cost of low-level programming and hardware management. Frameworks like Alchemist enable Spark to offload computation to optimized HPC libraries via MPI, combining usability with performance[10]. HadoopMapReduce brings parallel processing to very large datasets with minimal code (just map and reduce functions), but its batch-oriented nature and disk I/O can be slow for complex analytics. Spark improves on Hadoop by enabling general directed acyclic graphs (DAGs) of operations in memory, trading some hardware efficiency for ease of use and speed on certain workloads.

In practice, these tools are often used in combination; for instance, large-scale data pipelines might use Spark for preprocessing and an MPI-based HPC code for detailed simulation, with results integrated downstream [9].

IV. CASE STUDIES: REAL-WORLD APPLICATIONS OF PARALLEL COMPUTING FRAMEWORKS

To illustrate how parallel computing frameworks are applied in real-world scenarios, we present several case studies encompassing both HPC simulations and big data analytics.

A. Scientific Simulation on HPC

Large-scale scientific problems—such as climate modeling, astrophysics simulations, and computational genomics—rely on HPC clusters and supercomputers. These applications use MPI (often combined with OpenMP) to run simulations on thousands of CPU cores or GPUs in parallel [2]. For example, state-of-the-art climate models decompose the Earth's atmosphere and oceans into spatial grids and distribute different portions of the grid to different processors. Each processor computes the physics for its region and exchanges boundary data with neighbors at each time step using MPI messaging. This approach was employed in the Parallel Climate Model and its

successors, enabling more detailed and faster climate simulations as more processors are added.

Another example is molecular dynamics in biochemistry, where millions of atoms are simulated; a program like NAMD uses spatial decomposition with MPI, and can also offload force calculations to GPUs via CUDA for additional speedup. The general outcome is that problems which would take weeks or months to run on a single computer can be completed in hours or less on a parallel HPC system. Moreover, some simulations are only feasible on parallel systems—for instance, resolving fine-grained features in a supernova explosion or in aerodynamic simulations of an aircraft requires an exascale-level computational effort.

Modern supercomputers such as Frontier (USA) or Fugaku (Japan) provide this capability, combining tens of thousands of CPU cores and GPUs. HPC centers also emphasize high-throughput I/O and memory bandwidth to support data-intensive simulations (for example, writing simulation outputs or checkpoint files in parallel to a cluster file system).

B. Big Data Analytics with Spark

On the big data side, organizations routinely analyze datasets of terabyte or peta byte scale for insights. Apache Spark has been adopted as a unified engine for such analytics due to its speed and ease of use. A notable demonstration of Spark's capability was its record in the large-scale sort benchmark: Spark was able to sort 100 TB of data three times faster than Hadoop MapReduce while using only one-tenth of the computing resources.

This benchmark highlights how an in-memory, optimized framework can significantly outperform earlier paradigms on certain tasks. In practice, companies use Spark for tasks like processing event logs, running machine learning algorithms on large feature data sets, and interactive data exploration. For example, a social media company might use Spark to parallelize the computation of user engagement metrics or to train a recommendation system model on billions of data points. Spark's RDD-based approach ensures that if some servers in the cluster fail during computation, the system can recompute lost data partitions using lineage information, thereby completing the job without restarting the entire computation [11],[12]. Another aspect is the rich ecosystem: a data scientist can use PySpark (the Python API for Spark) to write a scalable data processing script, benefiting from parallelism under the hood without needing to explicitly manage threads or MPI

processes.

C. Big Data Meets HPC: SKA Telescope

The Square Kilometre Array (SKA) is a compelling example that blurs the line between HPC and big data analytics. The SKA, once operational, will continuously generate enormous streams of radio telescope data (on the order of petabytes per day) that need to be processed to produce scientific images and detect celestial signals. The data processing pipeline for SKA involves both HPC-style numerical algorithms (Fourier transforms for image synthesis, filtering, etc.) and big data management (distributed storage and streaming). The computing infrastructure for SKA is essentially an HPC cluster with tens of thousands of processors, yet it must handle data in a fault-tolerant way similar to big data systems because the data volume is so high and continuous [13].

Projects working on SKA data processing use workflow managers and custom frameworks (e.g., DALiUGe) to distribute tasks across the HPC resources. This case study underlines the convergence of techniques: the challenge is to achieve real-time parallel processing on an exascale data stream, requiring HPC algorithms to be extremely optimized and to run in parallel, and big data frameworks to provide resilience and easy reconfiguration in the face of failures or changing data rates [14].

D. Hybrid Analytics: Combining HPC and Big Data

In many real scenarios, HPC computations and big data analytics are part of a single end-to-end workflow. Consider an environmental science project where an HPC simulation generates a large dataset (e.g., a high-resolution climate projection). Once generated, this output might be tera bytes in size and needs to be analyzed—for example, to detect patterns or extreme events. A Spark or Hadoop cluster can be used to post-process the simulation output in parallel (for instance, extracting statistics or performing data mining across the dataset).

Conversely, machine learning models trained on big data are increasingly used to guide or accelerate HPC simulations (surrogate models). An example is using neural networks to approximate expensive physics calculations; training such networks might use GPUs in a data-parallel fashion on large simulation datasets, and then the inference (prediction) can be integrated into an MPI simulation to reduce its runtime. These hybrid workflows illustrate how parallel computing in data

processing is not confined to one paradigm—the best solution often combines HPC and big data techniques.

E. Recent Advancements

The landscape of parallel computing frameworks is continuously evolving to address the growing demands of scalability, efficiency, and flexibility. Several recent advancements reflect the trend toward integrating high-performance computing with cutting-edge AI techniques and big data pipelines.

One notable contribution is Deep RC[15], which introduces a scalable and efficient pipeline tailored for data engineering and deep neural network (DNN) training. This framework is designed to handle the complexities of large-scale data processing by combining parallel computing strategies with optimized DNN workflows. Deep RC addresses key bottlenecks in data preprocessing, model training, and deployment, enabling faster turnaround times and seamless scalability across distributed environments.

Another significant development is TaPS (Task-based Performance Suite) [16]. TaPS provides a comprehensive evaluation platform for assessing task-based execution frameworks. By systematically measuring performance, scalability, and reliability across various parallel environments, TaPS aids researchers and practitioners in selecting and fine-tuning the most suitable execution strategies for their work loads. Its contributions are particularly relevant as task-based programming models gain traction for their ability to exploit fine-grained parallelism while maintaining ease of use.

Additionally, ENRIQ offers an innovative approach to enterprise-scale data retrieval and querying. By leveraging neural network-based retrieval mechanisms, ENRIQ bridges the gap between traditional database systems and modern AI-driven search, ensuring fast and intelligent access to massive datasets. This system exemplifies how parallelism and AI techniques can be used to enhance data processing pipelines in business environments.

Complementing these frameworks, recent soft computing approaches underscore the integration of fuzzy logic, evolutionary algorithms, and machine learning techniques within big data contexts. These hybrid methods leverage the inherent parallelism of soft computing paradigms to process large, uncertain, and dynamic datasets efficiently, paving the way for more adaptable and intelligent data-driven solutions.

Collectively, these advancements demonstrate how parallel computing frameworks

are rapidly adapting to meet the challenges of modern data processing, blending traditional high-performance computing principles with AI and big data innovations.

V. DISCUSSION

The case studies and framework reviews above reveal a spectrum of parallel computing approaches. Here we discuss some comparative insights, challenges, and emerging trends.

A. Performance vs. Productivity

HPC frameworks like MPI and OpenMP generally deliver the highest performance for tightly-coupled computations, but require significant expertise to use effectively. The programmer must handle low-level details: data distribution, synchronization, load balancing, etc. Big data frameworks like Spark and Hadoop, on the other hand, sacrifice some efficiency in favor of developer productivity and fault tolerance. As noted earlier, MPI on InfiniBand networks can achieve extremely low-latency communication suitable for fine-grained parallel algorithms, whereas Spark's overhead would be too high for such workloads.

However, for coarse-grained tasks on huge datasets, Spark can provide at its factory performance with far less programmer effort, automatically handling data locality and failures. It has been observed that "MPI (with OpenMP) has mostly met users' needs for high computational performance, while big data frameworks such as Spark provide high-level programming, resiliency and I/O handling." The choice of framework thus depends on the problem: climate simulation code might need MPI to run efficiently on a supercomputer, while a log analytics task would use Spark to quickly get results on a cluster of commodity servers.

B. Fault Tolerance and Reliability

One of the stark differences between traditional HPC and big data processing is how they handle faults. HPC jobs typically run on dedicated hardware where failures, while possible, are not frequent for shorter runs; the software largely assumes reliable execution, with checkpoints as a safety net. Big data frameworks operate on large clusters (often with commodity or cloud nodes) where failures are expected. They incorporate fault tolerance into the runtime: for example, Hadoop replicates HDFS data blocks across nodes so that a node loss does not lose data, and Spark's RDD lineage enables recomputation of lost partitions.

This makes big data jobs more robust to node churn and suitable for long-running services (like continuously updating computations in streaming analytics). As HPC moves towards exascale, however, the probability of component failure during a job increases (simply due to scale). Techniques from the big data world are influencing HPC: task-based programming models in HPC (like Legion or OpenMP tasking) can potentially restart just failed tasks, and there is research into algorithm-based fault tolerance. Conversely, big data systems are beginning to leverage HPC ideas to reduce overhead—such as using RDMA networks for faster data shuffle, or specialized file systems to speed up I/O.

C. Heterogeneous Computing Challenges

The rise of GPUs and other accelerators is a double-edged sword. They offer massive performance, but they also complicate software. HPC developers have to manage multiple programming environments (CPU code, GPU kernels) and data movement between them. This has led to frameworks like OpenACC and enhancements in OpenMP to offload computations to GPUs. While these make accelerator usage easier, achieving good performance still often requires algorithm redesign to fit GPU architectures.

In the big data domain, initial frameworks were CPU-centric (e.g., Hadoop and Spark originally did not utilize GPUs). Efforts like Hetero Spark and GPU Spark show an appetite for incorporating GPUs to accelerate big data processing (for example, accelerating machine learning tasks within a Spark pipeline) [17]. A challenge is that big data frameworks use high-level languages (Java/Scala, Python) that are not naturally suited to low-level GPU programming. Bridging this gap requires careful integration (e.g., using JNI to call CUDA code, or relying on GPU-aware libraries). There is ongoing work to create unified programming models that can schedule tasks across CPU and GPU resources seamlessly. The overall trend is that heterogeneity is increasing, with not just GPUs but also FPGAs and AI accelerators in systems. The programming models and runtime systems must evolve to schedule and optimize across these diverse resources, a non-trivial task that is a current research focus.

D. Convergence of HPC and Big Data

As data analytics tasks grow more computationally demanding and HPC simulations produce ever-larger data outputs, the distinction

between HPC and big data computing is blurring. The concept of HPC-big data convergence is evident in initiatives like the Big Data and Exascale Computing (BDEC) series, which aim to find common ground between the two ecosystems.

One aspect of convergence is in the software stack: for example, using containerization and cloud orchestration tools to deploy HPC applications, or using HPC-style scheduling for big data workloads to improve utilization on large clusters. Another aspect is unified frameworks: there is interest in frameworks that can handle both data flow (stream or batch data processing) and compute-heavy tasks. Projects such as TensorFlow for machine learning also embody a mix—TensorFlow uses dataflow graphs (like a big data system) but executes many linear algebra operations on GPUs (like HPC), often using MPI for multi-node communication in its high-performance implementations.

Our review reveals that MPI and OpenMP offer superior control for tightly coupled HPC tasks, whereas Spark and Hadoop provide ease of programming and resilience. Emerging frameworks like Merak[8] and GraphTensor continue to bridge HPC and big data needs. Recent soft computing applications also leverage parallelism for scalability.

Looking forward, researchers are considering paradigms that combine the best of both worlds: the rich semantics and fault tolerance of big data frameworks with the efficiency and low-level control of HPC frameworks. One idea is having HPC applications output their data into in-memory data stores that Spark or Flink can immediately pick up for analysis, eliminating slow file I/O. Another is using Apache Arrow or similar memory formats to allow zero-copy data exchange between HPC codes (in C++) and big data systems (in Java/Python). On the hardware side, future exascale systems will likely run complex workflows that include data analytics, requiring job schedulers to handle diverse MPI and Spark jobs, possibly with dynamic resource sharing.

In summary, there is a clear trade-off in current parallel computing approaches: HPC techniques offer maximum performance and scalability for tightly coupled computations, while big data techniques offer ease of programming and resilience for loosely coupled, data-intensive workloads [18]. The state-of-the-art often involves combining approaches to solve end-to-end problems [19]. The challenges of heterogeneity and scalability are driving the two communities together, and we expect continued innovation in frameworks that can span the requirements of both

HPC and big data processing. Nextflow addresses reproducibility challenges in parallel data processing pipelines by supporting portable, container-based execution [20].

VI. CONCLUSION

Parallel computing is an essential enable for modern data processing challenges, from simulating complex physical phenomena to extracting insights from massive datasets. This paper has surveyed the landscape of parallel computing frameworks, highlighting how each addresses the twin goals of accelerating computation and managing large-scale data. Shared-memory and distributed-memory architectures provide the hardware foundation for parallelism, while programming models like MPI, OpenMP, CUDA, MapReduce, and Spark realize parallelism in software.

Our review has underscored that no single model is optimal for all tasks: rather, the choice depends on the problem characteristics and performance requirements. High-performance computing frameworks excel at squeezing out computational performance on specialized hardware, and big data frameworks shine in productivity and fault tolerance on large clusters. In practice, hybrid approaches are increasingly common, leveraging the strengths of each. Real-world case studies demonstrate impressive achievements, from exascale simulations to real-time big data analytics, all powered by parallel computing techniques.

As hardware continues to evolve towards greater core counts and specialized accelerators, parallel computing frameworks will need to adapt, emphasizing interoperability and ease of use without sacrificing performance. The ongoing convergence of HPC and big data paradigms points toward a future computing ecosystem where scientific computing and data analytics seamlessly integrate. By understanding the current frameworks and their trade-offs, researchers and practitioners can better design systems and applications that harness parallel computing effectively for their data processing needs. The continual innovations in this field promise to further expand the boundaries of what can be computed and analyzed, driving discovery and decision-making in the data-driven world.

REFERENCES

- [1] S.Wei, F. Wang, H.Deng, C.Liu, W.Dai, B.Liang, Y.Mei, C. Shi, Y. Liu, and J. Wu, "Open cluster: A flexible distributed computing framework for astronomical data

- processing,” arXiv preprint arXiv:1701.04907, 2017. [Online]. Available: <https://arxiv.org/abs/1701.04907>
- [2] S. Habib, V. Morozov, and H. Finkel, “Hacc: Extreme scaling and performance across diverse architectures,” in SC’13: Proceedings of the International Conference on High Performance Computing, Networking, Storage and Analysis, 2013. [Online]. Available: <https://dl.acm.org/doi/10.1145/2503210.2504566>
- [3] A. Eichenberger, J. Mellor-Crummey, and M. Schulz, “Openmp application programming interface v5.1,” 2020. [Online]. Available: <https://www.openmp.org/wp-content/uploads/ompt-tr2.pdf>
- [4] S. K. Muller and J. Hoffmann, “Modeling and analyzing evaluation cost of cuda kernels,” Proc. ACM Program. Lang., vol. 5, no. POPL, 2021. [Online]. Available: <https://doi.org/10.1145/3434306>
- [5] D. Chang, L. Li, Y. Chang, and Z. Qiao, “Implementation of map reduce parallel computing framework based on multi-data fusion sensors and gpu cluster,” EU RASIP Journal on Advances in Signal Processing, vol. 2021, no. 1, p. 77, 2021. [Online]. Available: <https://asp-eurasipjournals.springeropen.com/articles/10.1186/s13634-021-00787-7>
- [6] J. Dean and S. Ghemawat, “Mapreduce: Simplified data processing on large clusters,” Communications of the ACM, vol. 51, no. 1, pp. 107–113, 2008. [Online]. Available: <https://doi.org/10.1145/1327452.1327492>
- [7] K. Shvachko, H. Kuang, S. Radia, and R. Chansler, “The hadoop distributed file system,” Proceedings of the 2010 IEEE 26th Symposium on Mass Storage Systems and Technologies (MSST), pp. 1–10, 2010. [Online]. Available: <https://doi.org/10.1109/MSST.2010.5496972>
- [8] Z. Lai, S. Li, X. Tang, K. Ge, W. Liu, Y. Duan, L. Qiao, and D. Li, “Merak: An efficient distributed dnn training framework with automated 3d parallelism for giant foundation models,” arXiv preprint arXiv:2206.04959, 2022. [Online]. Available: <https://arxiv.org/abs/2206.04959>
- [9] J. Jang, M. Kwon, D. Gouk, H. Bae, and M. Jung, “Graphtensor: Comprehensive gnn-acceleration framework for efficient parallel processing of massive datasets,” arXiv preprint arXiv:2305.17469, 2023. [Online]. Available: <https://arxiv.org/abs/2305.17469>
- [10] A. Gittens, K. Rothauge, S. Wang, M. W. Mahoney, L. Gerhardt, Prabhat, J. Kottalam, M. Ringenburt, and K. Maschhoff, “Accelerating large-scale data analysis by offloading to high-performance computing libraries using alchemist,” arXiv preprint arXiv:1805.11800, 2018. [Online]. Available: <https://arxiv.org/abs/1805.11800>
- [11] M. Zaharia, M. Chowdhury, and T. Das, “Resilient distributed data sets: A fault-tolerant abstraction for in-memory cluster computing,” in Proc. 9th USENIX Symposium on Networked Systems Design and Implementation (NSDI), 2012, pp. 15–28. [Online]. Available: <https://www.usenix.org/conference/nsdi12/technical-sessions/presentation/zaharia>
- [12] M. Zaharia, M. Chowdhury, M. J. Franklin, S. Shenker, and I. Stoica, “Spark: Cluster computing with working sets,” Proceedings of the 2nd USENIX Conference on Hot Topics in Cloud Computing, 2010. [Online]. Available: https://www.usenix.org/legacy/event/hotcloud10/tech/full_papers/Zaharia.pdf
- [13] R. Wang, R. Tobar, M. Dolensky, T. An, A. Wicenec, C. Wu, Dulwich, N. Podhorszki, V. Anantharaj, E. Suchyta, B. Lao, and S. Klasky, “Processing full-scale square kilometre array data on the summit super computer,” in Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis. IEEE Press, 2020. [Online]. Available: <https://dl.acm.org/doi/10.5555/3433701.3433704>
- [14] C. Wu, R. Tobar, K. Vinsen, A. Wicenec, D. Pallo, B. Lao, R. Wang, T. An, M. Boulton, I. Cooper, R. Dodson, M. Dolensky, Y. Mei, and F. Wang, “Daliuge: A graph execution framework for harnessing the astronomical data deluge,” Astronomy and Computing, vol. 20, pp. 1–15, 2017. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2213133716301214>
- [15] H. Zhao, Z. Yang, Y. Cheng, C. Tian, S. Ren, W. Xiao, M. Yuan, L. Chen, K. Liu, Y. Zhang, Y. Li, and W. Lin, “Goldminer: Elastic scaling of training data pre-processing pipelines for deep learning,” Proc. ACM Manag. Data, vol. 1, no. 2, Jun. 2023. [Online]. Available: <https://doi.org/10.1145/3589773>
- [15] T. Grass, A. Rico, M. Casas, M. Moreto, and E. Ayguade, “Task point: Sampled

- simulation of task-based programs,” in 2016 IEEE International Symposium on Performance Analysis of Systems and Software (ISPASS). IEEE, 2016, pp. 296–306.[Online]. Available: https://ieeexplore.ieee.org/abstract/document/7482104?casa_token=PP9KPsDDdEsAAAAA:YxoFNKGP610Gxs86ZU4uWVmKfh3nV4nm2oUuB3xSJUJvtw01HbTwpwaGhofP60SHvqVkrZg
- [16] C. Wu, R. Tobar, K. Vinsen, A. Wicenc, D. Pallo t, B. Lao, R. Wang, T. An, M. Boulton, I. Cooper et al., “Daliuge: A graph execution framework for harnessing the astronomical data deluge,” arXiv preprint arXiv:1702.07617, 2017. [Online]. Available: <https://arxiv.org/abs/1702.07617>
- [17] B. Carver, J. Zhang, A. Wang, A. Anwar, P. Wu, and Y. Cheng, “Wukong: A scalable and locality-enhanced framework for server less parallel computing,” arXiv preprint arXiv:2010.07268, 2020.[Online]. Available: <https://arxiv.org/abs/2010.07268>
- [18] M. Abadi, “Tensorflow: A system for large-scale machine learning,” 12th USENIX Symposium on Operating Systems Design and Implementation (OSDI16), pp. 265–283, 2016.[Online]. Available: <https://www.usenix.org/system/files/conference/osdi16/osdi16-abadi.pdf>
- [19] P. Di Tommaso, M. Chatzou, E. W. Floden, P. P. Barja, E. Palumbo, and C. Notredame, “Nextflow enables reproducible computational workflows,” Nature biotechnology, vol. 35, no. 4, pp. 316–319, 2017.[Online]. Available: <https://www.nature.com/articles/nbt.3820>