

Big Data: Analysis of Large and Complex Data Sets

Syed Faheemuddin*

Osmania University, Hyderabad, Telangana, India

*Corresponding author: Syed Faheemuddin, s.faheem@alramzre.com

Abstract

The paradigm of Big Data represents a fundamental shift in the scale and complexity of information available for analysis, characterized by the foundational Four V's: immense Volume, rapid Velocity, diverse Variety, and uncertain Veracity. This paper provides a comprehensive analysis of the methodologies and technologies required to process and extract value from these large and complex datasets. It examines the transition from traditional relational databases to distributed computing frameworks, such as Hadoop and Spark, which enable the storage and parallel processing of data across clusters of commodity hardware. The discussion extends to the ecosystem of NoSQL databases, designed to handle the heterogeneity of unstructured and semi-structured data. The analytical spectrum is explored, from descriptive and diagnostic analytics to the more advanced applications of predictive and prescriptive modeling, heavily reliant on machine learning and statistical algorithms to uncover patterns and generate actionable insights. However, this pursuit is not without significant challenges. This article critically addresses impediments including data quality and preprocessing, substantial infrastructure demands, a pronounced skills gap, and profound ethical concerns regarding data privacy, security, and algorithmic bias. Ultimately, this analysis concludes that effectively navigating the big data labyrinth necessitates a sophisticated integration of technology, analytical expertise, and robust ethical governance to fully harness its transformative potential for innovation and decision-making across scientific, commercial, and social domains.

Keywords

Big Data, Data Analytics, Hadoop, Spark, Machine Learning, Artificial Intelligence, IoT, Cloud Computing, Data Governance, Quantum Computing

1. Introduction

The exponential growth of data in recent decades has transformed the way individuals, businesses, and governments operate. The concept of Big Data refers to massive and complex datasets that cannot be efficiently captured, stored, managed, or analyzed using traditional data management tools [1]. Big Data is characterized by its volume, velocity, variety, veracity, and value, making it a critical asset in the digital age. These features highlight not only the size of the data but also its speed of generation, diversity of formats, reliability, and potential to generate insights [2]. The proliferation of smartphones connected devices, sensors, cloud computing, and social media platforms has created an ever-expanding stream of structured, semi-structured, and unstructured data. According to industry reports, global data is expected to reach 175 zettabytes by 2025, emphasizing the urgent need for advanced analytical techniques [3]. Analyzing Big Data enables organizations to extract hidden patterns, trends, and correlations that support decision-making, optimize operations, and create innovative solutions across various domains, including healthcare, finance, manufacturing, education, and government [4]. As a result, Big Data is not just a technological phenomenon but a strategic resource, reshaping economies, driving innovation, and influencing policymaking on a global scale [5].

2. Literature Review

The field of Big Data has evolved rapidly, with academic and industry research emphasizing its significance in modern analytics. Katal, A., Wazid, M., & Goudar, R. H. [6] first introduced the concept of the "3Vs" of Big Data-volume, velocity, and variety-later expanded to include veracity and value, forming the widely recognized 5Vs framework. Studies have shown that the increasing complexity of data requires distributed systems such as Hadoop and Apache Spark (Sagiroglu, S., & Sinanc, D.) [7] which enable organizations to process and analyze massive datasets efficiently.

Recent work has focused on machine learning-driven analytics, integrating artificial intelligence (AI) with Big Data pipelines to deliver real-time decision-making capabilities (Gandomi, A., & Haider, M.) [8]. Additionally, research highlights the importance of data lakes and cloud-based platforms, which provide scalable storage and computing resources (Demchenko, Y., Grosso, P., De Laat, C., & Membrey, P)[9]. Security and privacy challenges, however, remain a key concern, with studies emphasizing the need for data governance frameworks and ethical AI (Hilbert, M.) [10].

The literature reveals three primary areas of focus:

2.1 Technological Evolution

Development of distributed frameworks like Hadoop, Spark, and Kubernetes for managing large-scale analytics.

2.2 Industrial Applications

Adoption of Big Data analytics across domains such as precision medicine, fraud detection, smart cities, and recommendation systems.

2.3 Emerging Trends

Integration of edge computing, blockchain, and quantum computing into data processing ecosystems.

By reviewing these studies, this paper builds on established knowledge and identifies gaps, particularly in data governance strategies and the scalability of analytics in the era of AI-driven decision-making.

3. Defining Big Data and the 5Vs Framework

Big Data refers to datasets that are too large, complex, and fast-moving to be effectively captured, stored, and analyzed using traditional database systems. Unlike conventional data, which is typically structured and stored in relational databases, Big Data encompasses a wide range of data types, including structured, semi-structured, and unstructured data originating from multiple sources such as social networks, IoT devices, e-commerce platforms, and enterprise systems. The definition of Big Data extends beyond sheer size; it emphasizes the technological, analytical, and organizational challenges associated with handling and extracting value from massive datasets (McAfee, A., & Brynjolfsson, E) [11].

To conceptualize these challenges, analysts and researchers often use the 5Vs framework (Reinsel, D., Gantz, J., & Rydning, J) [12] which characterizes Big Data along five dimensions:

The complexity and strategic importance of Big Data can be better understood through the 5Vs framework (Figure 1), which captures its key dimensions: volume, velocity, variety, veracity, and value.

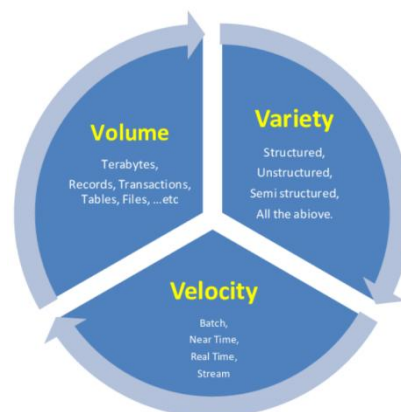


Figure 1. Framework

3.1 Volume

Volume refers to the sheer size of datasets, which are often measured in terabytes, petabytes, or even exabytes. As of 2025, global data creation is expected to exceed 175 zettabytes (Trigyn) [13]. Such scale requires scalable storage solutions like data lakes, distributed file systems (HDFS), and cloud storage services (Amazon S3, Azure Blob).

3.2 Velocity

Velocity describes the speed at which data is generated, processed, and analyzed. Modern applications such as stock trading platforms, connected vehicles, and smart manufacturing generate massive data streams that require real-time or near-real-time analytics. Tools such as Apache Kafka and Apache Flink are widely used to manage high-speed data ingestion and streaming analytics.

3.3 Variety

Variety captures the diverse forms of data collected from multiple sources. Unlike traditional datasets, which are typically structured (rows and columns), Big Data includes:

Structured data: Relational databases, transactional records.

Semi-structured data: JSON, XML logs.

Unstructured data: Images, videos, social media text, sensor data.

Managing and integrating these diverse formats requires flexible data models and processing frameworks like NoSQL databases (MongoDB, Cassandra).

3.4 Veracity

Veracity emphasizes the uncertainty and quality of data. With increasing data sources, organizations must address data inconsistencies, duplication, noise, and biases to ensure accurate analytics. Data cleansing techniques, metadata management, and governance frameworks are critical for improving veracity.

3.5 Value

Value is the ultimate goal of Big Data analytics: turning raw data into actionable insights. High-value applications include predictive maintenance, fraud detection, customer segmentation, and personalized recommendations. Organizations investing in Big Data infrastructure focus heavily on value extraction, as it directly impacts strategic decision-making and revenue growth.

The 5Vs model illustrates that Big Data is not merely about size but also about complexity, diversity, and strategic importance. Addressing these dimensions requires robust architectures, distributed processing frameworks, and analytical tools that support large-scale, real-time decision-making.

4. Sources of Big Data

Big Data originates from a variety of sources that contribute to its diversity, volume, and complexity. Understanding these sources is critical for designing efficient data pipelines and analytics systems. Key sources include:

The main categories of Big Data sources are summarized and visualized in Figure 2, illustrating the wide diversity of input streams that drive modern analytics.



Figure 2. Bigdata.

4.1 Social Media Platforms

Social networks such as Facebook, Instagram, Twitter (X), and TikTok generate massive amounts of unstructured data in the form of text, images, and videos. Organizations analyze this data for sentiment analysis, marketing campaigns, and user engagement strategies.

4.2 Internet of Things (IoT) Devices

Billions of connected devices, ranging from smart home appliances to industrial sensors, produce continuous streams of real-time data. IoT-generated Big Data is essential in predictive maintenance, smart cities, and supply chain optimization (Lee & Lee) [2].

4.3 Enterprise Systems and Transactions

Organizations collect structured data from ERP systems, CRM platforms, e-commerce portals, and financial transactions. This data is crucial for decision-making, forecasting, and fraud detection.

4.4 Scientific Research and Experiments

Fields like genomics, astronomy, and particle physics generate extremely large datasets through high-throughput sequencing, telescopes, and experiments such as those at CERN. These require advanced HPC (High-Performance Computing) systems for processing.

4.5 Machine-Generated Logs

Data from server logs, network monitoring tools, and cybersecurity systems provide insights for IT infrastructure optimization and threat detection.

4.6 Multimedia Content

Videos, music, images, and streaming data contribute significantly to Big Data growth, especially from platforms like YouTube, Netflix, and Spotify.

4.7 Open Data and Government Records

Government agencies publish open datasets for public policy, transportation, and urban planning. These datasets are often leveraged by researchers and organizations for analytics.

5. Big Data Architecture

Big Data Architecture is a comprehensive framework that defines the components, processes, and technologies used to ingest, process, store, manage, and analyze data that is too large or complex for traditional database systems. It is designed to be scalable, fault-tolerant, and flexible to handle the Volume, Velocity, and Variety of data.

5.1 Core Objectives of the Architecture

Scalability: Ability to handle growing data volumes by adding more nodes to the system (horizontal scaling).

Fault Tolerance: Ability to withstand hardware or software failures without losing data or interrupting processing.

Flexibility: Support for various data types (structured, semi-structured, unstructured) and multiple processing paradigms (batch, real-time).

Decoupling: Components are designed to be independent, allowing them to be scaled and managed separately. This is often achieved using message queues.

5.2 The Layered Architecture of a Big Data System

A robust big data pipeline is typically structured in layers, each with a specific responsibility.

Data Ingestion Layer

This layer is responsible for collecting data from diverse sources and bringing it into the system.

Tools

Apache Kafka: The de facto standard for a distributed, highly scalable, and fault-tolerant publish-subscribe messaging system. It acts as a durable buffer for high-velocity data streams.

Apache NiFi: Provides a user-friendly interface to design data flows (called dataflows) for automating data movement from various sources.

AWS Kinesis / Google Cloud PubSub: Managed cloud services for ingesting and processing real-time data streams.

Function: Decouples data producers from consumers, ensuring no data is lost during peak loads and allowing downstream systems to process data at their own pace.

Data Storage Layer

This layer stores the massive volumes of ingested data in its raw or processed form.

Data Lakes: A central repository that stores structured and unstructured data *in its raw format*. This allows for schema-on-read, meaning the structure is applied when the data is read, not when it's written, providing immense flexibility.

Technologies: Often built on top of HDFS (Hadoop Distributed File System) or, more commonly now, cloud object storage like Amazon S3, Azure ADLS, or Google Cloud Storage due to its durability, scalability, and cost-effectiveness.

Distributed File Systems: HDFS is designed to run on clusters of commodity hardware, storing large files across multiple machines.

NoSQL Databases: Used for specific applications requiring low-latency access.

Key-Value Stores (e.g., Redis, DynamoDB): For simple, fast queries based on a key.

Document Stores (e.g., MongoDB, Couchbase): For storing semi-structured JSON-like documents.

Columnar Stores (e.g., Apache Cassandra, HBase): Optimized for reading and writing large volumes of data by columns, not rows.

Graph Databases (e.g., Neo4j): For storing and querying complex relationships.

5.3 Data Processing Layer

This is the computational engine of the architecture, where data is transformed, cleaned, aggregated, and prepared for analysis.

Batch Processing: Processing large, finite datasets in jobs that run over a period of time (e.g., hours).

Technologies: Apache Hadoop MapReduce (the original paradigm), Apache Spark (faster due to in-memory processing). Used for ETL (Extract, Transform, Load) jobs, daily reports, and model training.

Stream Processing: Processing unbounded data streams in real-time or near-real-time as they arrive.

Technologies: Apache Spark Streaming (micro-batching), Apache Flink (true streaming with low latency), Apache Samza, and Kafka Streams. Used for fraud detection, real-time recommendations, and live dashboards.

Hybrid Processing (Lambda/Kappa):

Lambda Architecture: Runs both batch and stream processing pipelines in parallel, merging their results to provide a complete view. This is complex to maintain.

Kappa Architecture: Simplifies the pipeline by using a single stream processing engine for all data. To reprocess data, you simply replay the historical data stream.

The Lambda approach is illustrated in Figure 3, which contrasts batch and stream processing pipelines to demonstrate their complementary roles within Big Data systems.

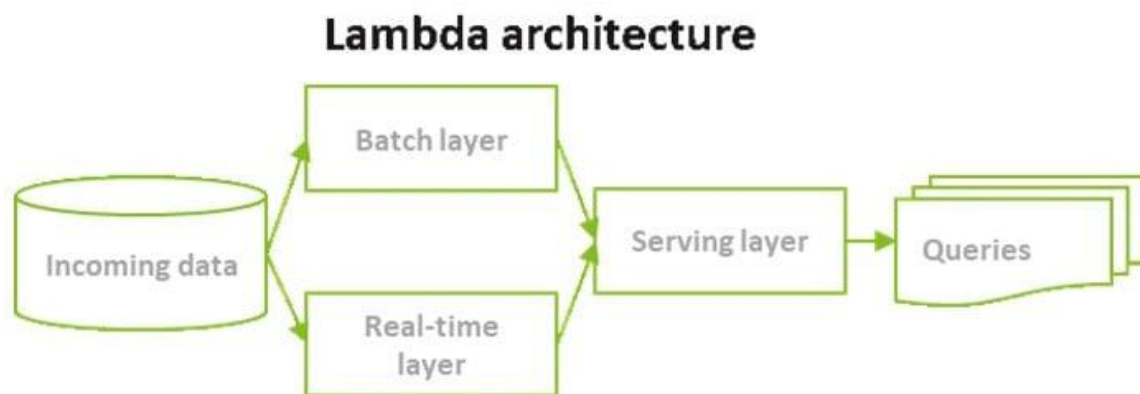


Figure 3. Lambda Architecture

5.4 Data Analysis & Serving Layer

This layer provides interfaces for users, applications, and data scientists to access the processed data.

Data Warehouses: Store processed, structured data in a schema optimized for complex SQL queries and business intelligence (BI).

Technologies: Amazon Redshift, Google BigQuery, Snowflake, Azure Synapse Analytics. These are often used as the "serving layer" for BI tools.

Analytical Engines: Apache Hive (SQL-on-Hadoop), Presto/Trino (fast distributed SQL query engine for data lakes), Apache Drill.

Machine Learning & AI Tools: Jupyter Notebooks, MLflow, Kubeflow integrated with frameworks like TensorFlow and PyTorch to build models on the stored data.

5.5 Data Orchestration & Management Layer

The "central nervous system" that coordinates and manages the entire data pipeline.

Workflow Schedulers: Apache Airflow (the industry standard), Luigi, Apache Oozie. They define, schedule, and monitor complex data workflows as Directed Acyclic Graphs (DAGs).

Metadata Management & Data Cataloging: Apache Atlas, Amundsen, DataHub. These tools track the lineage of data (where it came from, how it was transformed), manage data glossary, and provide a searchable catalog of all data assets.

Cluster Resource Management: Apache YARN (Yet Another Resource Negotiator) and Kubernetes. They manage and schedule computational resources across the cluster, deciding which jobs get which resources (CPU, memory).

A holistic view of Big Data system components and their layered interactions is presented in Figure 4, offering a structured perspective on the architecture.

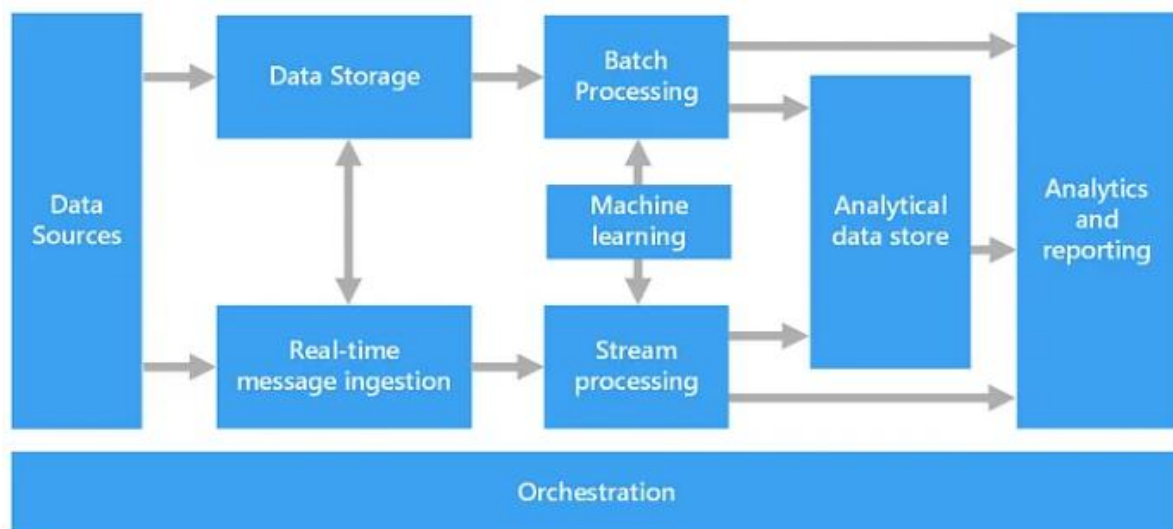


Figure 4. Data Architecture

6. Big Data Tools and Technologies

The big data ecosystem is vast, comprising open-source projects, commercial offerings, and cloud-native services. These tools are designed to handle the challenges of Volume, Velocity, and Variety. They can be best understood by mapping them to the layers of a modern data architecture.

6.1 Data Ingestion & Collection Tools

These tools are responsible for gathering data from various sources and transporting it into the data system.

Apache Kafka: The industry standard for a distributed event streaming platform. It is not just a message queue; it is used for building real-time data pipelines and streaming applications. It excels at handling high-throughput, fault-tolerant data ingestion.

Apache NiFi: An easy-to-use, powerful, and reliable system to process and distribute data. It provides a web-based UI to design data flows for automation and supports a vast array of sources and destinations.

Fluentd: An open-source data collector for unified logging layers. It allows you to unify data collection and consumption for better use and understanding of data.

6.2 Cloud Services

AWS Kinesis: Suite of services (Kinesis Data Streams, Firehose) for collecting, processing, and analyzing real-time streaming data.

Google Cloud Pub/Sub: A scalable event ingestion service for streaming analytics and data integration.

Azure Event Hubs: A big data streaming platform and event ingestion service.

6.3 Data Storage & Management Tools (The Data Lake/Warehouse)

These technologies store the massive volumes of data in a scalable and cost-effective manner.

6.3.1 Distributed File Systems & Object Storage

Hadoop Distributed File System (HDFS): The original scalable storage system for Hadoop, designed to run on commodity hardware.

Cloud Object Storage: The modern default for data lakes due to infinite scalability, durability, and cost. Examples include Amazon S3, Azure Blob Storage, and Google Cloud Storage.

6.3.2 NoSQL Databases (Handling Variety)

Document Databases: Store data in JSON-like documents (e.g., MongoDB, Couchbase).

Columnar Databases: Optimized for queries over large datasets by storing data in columns instead of rows (e.g., Apache Cassandra, HBase, ScyllaDB).

Key-Value Stores: Simple model using key-value pairs for high-performance lookups (e.g., Redis (in-memory), Amazon DynamoDB).

Graph Databases: Store data focused on relationships and connections (e.g., Neo4j, Amazon Neptune).

Figure 5 provides a comprehensive overview of the Big Data ecosystem, mapping the major technological domains—including data storage, processing frameworks, analytics tools, and visualization platforms—and illustrating how these components interact to enable end-to-end data-driven decision-making.

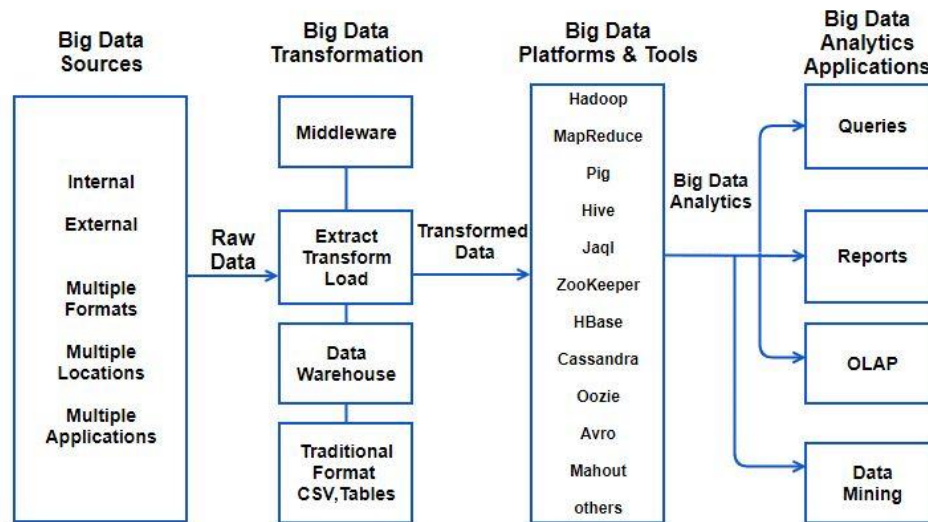


Figure 5. Big Data

6.3.3 Data Warehouses (Optimized for Analytics)

Snowflake: A cloud-native, SaaS-based data warehouse that separates storage and compute, offering high performance and concurrency.

Amazon Redshift: A fast, fully managed, petabyte-scale data warehouse service.

Google BigQuery: A serverless, highly scalable, and cost-effective multi-cloud data warehouse.

Apache Druid: A real-time analytics database designed for low-latency queries on event-based data.

6.4 Data Processing & Computation Frameworks

These are the engines that transform, clean, aggregate, and analyze the stored data.

6.4.1. Batch Processing (Processing large, finite datasets)

Apache Hadoop MapReduce: The original programming model for processing large datasets in parallel across a cluster.

Apache Spark: The dominant engine for large-scale data processing. It is faster than MapReduce due to its in-memory processing capabilities and supports SQL, streaming, machine learning, and graph processing in a unified framework.

6.4.2. Stream Processing (Processing infinite data in real-time)

Apache Flink: A high-performance, true streaming engine designed for low latency and high throughput. It handles event-time processing and state management elegantly.

Apache Spark Streaming: Uses micro-batching to provide streaming capabilities on top of the Spark core.

Apache Samza: A stream processing framework often used with Kafka for stateful processing.

Apache Storm: One of the first widely used stream processing frameworks.

6.4.3 Cloud-Native Processing

AWS Glue / Azure Data Factory / Google Cloud Dataflow: Managed ETL (Extract, Transform, Load) services that simplify the process of preparing and loading data for analysis. Dataflow is based on the Apache Beam model.

Several tools enable organizations to manage and analyze massive datasets:

Key frameworks and their categorization are presented in Figure 6, outlining major tools and technologies across storage, data ingestion, processing, analytics, and visualization, and highlighting their roles within the Big Data ecosystem.

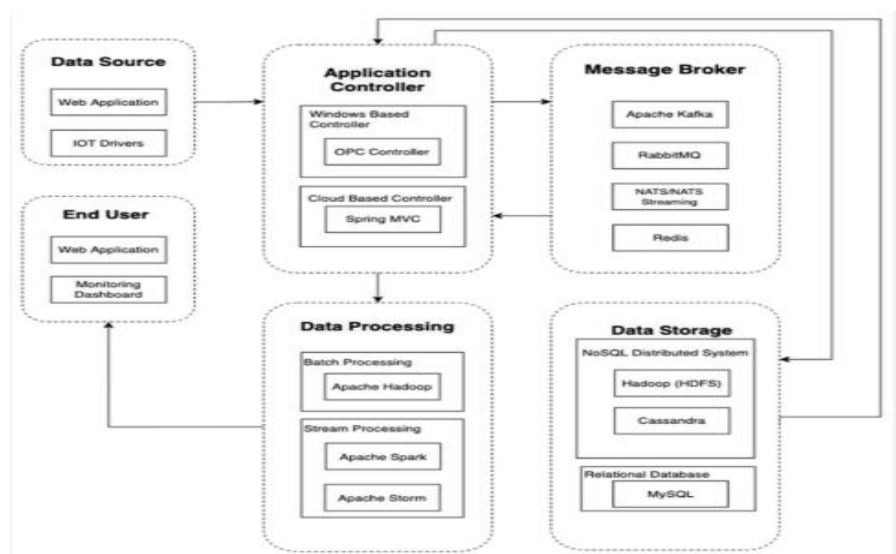


Figure 6. Big Data processing frameworks.

Category	Tools/Technologies	Purpose
Storage Systems	HDFS, Amazon S3, Google Cloud Storage	Distributed data storage
Processing Frameworks	Apache Hadoop, Apache Spark, Apache Flink	Batch & stream processing
NoSQL Databases	MongoDB, Cassandra, HBase	Scalable, schema-less storage
Data Ingestion	Apache Kafka, Apache NiFi, AWS Kinesis	Real-time data streaming
Analytics Platforms	TensorFlow, PyTorch, Scikit-learn	Machine learning & deep learning
Visualization Tools	Tableau, Power BI, D3.js	Visual representation of insights
Cloud Platforms	AWS, Azure, Google Cloud BigQuery	Scalable data management & analytics

7. Big Data Analytics Approaches

Big Data analytics can be classified into four primary approaches:

Descriptive Analytics: Summarizes historical data for trend analysis.

Diagnostic Analytics: Explains why events occurred using statistical modeling.

Predictive Analytics: Uses ML algorithms to forecast future outcomes.

Prescriptive Analytics: Recommends optimal strategies for decision-making.

Challenges in applying these approaches are highlighted in Figure 7, emphasizing key concerns such as scalability, privacy, security, and data governance within Big Data systems.

10 big data challenges

To capitalize on big data, companies will need to successfully master these challenges.



Figure 7. Data Challenges.

8. Applications of Big Data

Big Data analytics drives transformation across industries:

Healthcare: Precision medicine, disease outbreak prediction.

Finance: Fraud detection, algorithmic trading.

Retail: Customer segmentation, recommendation systems.

Manufacturing: Predictive maintenance, supply chain analytics.

Smart Cities: Traffic management, energy optimization.

Cybersecurity: Threat intelligence and real-time monitoring.

9. Challenges in Big Data Adoption

Despite its potential, Big Data faces several obstacles:

Data Privacy and Security: Regulatory compliance (GDPR, HIPAA).

Infrastructure Costs: High costs for scalable storage and computation.

Data Quality Issues: Noise, duplication, and biases.

Talent Shortage: Lack of skilled data scientists and engineers.

Interoperability: Integrating diverse tools and platforms.

10. Future Trends in Big Data

Emerging trends indicate that Big Data will become even more transformative:

Edge Computing: Processing data closer to devices to reduce latency.

Quantum Computing: Solving computationally complex problems at scale.

Blockchain for Data Integrity: Ensuring trust and transparency in data transactions.

AI-Powered Analytics: Automating insight generation through generative AI.

Data Democratization: Empowering non-technical users with self-service analytics tools.

11. Conclusion

Big Data has fundamentally revolutionized decision-making, transitioning it from a realm of intuition and limited sampling to one driven by comprehensive, evidence-based insights derived from vast and complex datasets. This paradigm shift empowers organizations across every sector-from healthcare and finance to manufacturing and urban planning-to optimize operations, personalize customer experiences, and innovate at an unprecedented pace.

As we look to the future, the convergence of Big Data with other transformative technologies will further amplify its impact. The integration of Artificial Intelligence (AI) and Machine Learning will move analytics beyond description and prediction into the realm of autonomous decision-making and continuous learning, creating self-optimizing systems. The proliferation of the Internet of Things (IoT) will exponentially increase the volume and velocity of data, providing a real-time digital pulse of the physical world, from global supply chains to individual health metrics.

However, this path forward is not without significant hurdles that must be proactively addressed:

Privacy and Security: As data collection becomes more pervasive, robust ethical frameworks and stringent data governance models are essential to protect individual privacy and prevent breaches. Regulations like GDPR are just the beginning.

Scalability and Infrastructure: The relentless growth in data volume demands ever-more efficient and cost-effective storage and processing solutions, pushing the limits of current cloud and distributed computing models.

Interoperability: The true value of data is unlocked when it can be seamlessly integrated and shared. Overcoming silos and establishing universal standards for data format and exchange remain critical challenges.

The next frontier of Big Data will be shaped by emerging technologies that promise to overcome these very challenges and unlock new possibilities:

In essence, the future of Big Data promises a move towards greater automation, efficiency, and global impact. It is a future where data doesn't just inform but acts-intelligently, securely, and instantly-driving progress and solving some of the world's most pressing challenges. The organizations that succeed will be those that not only harness these technologies but also master the ethical and strategic imperatives of responsible data use.

References

- [1] Reinsel, D., Gantz, J., & Rydning, J. (1). The Digitization of the World: From Edge to Core. IDC White Paper.
- [2] Marr, B. (2). Big Data in Practice: How 45 Successful Companies Used Big Data Analytics to Deliver Extraordinary Results. Wiley.
- [3] Manyika, J., Chui, M., Brown, B., Bughin, J., Dobbs, R., Roxburgh, C., & Hung Byers, A. (3). Big Data: The Next Frontier for Innovation, Competition, and Productivity. McKinsey Global Institute.
- [4] Goudar, I. A. T., Yaqoob, I., Anuar, N. B., Mokhtar, S., Gani, A., & Khan, S. U. (4). The rise of “big data” on cloud computing: Review and open research issues. *Information Systems*, 47, 98–115. <https://doi.org/10.1016/j.is.2014.07.006>
- [5] Chen, M., Mao, S., & Liu, Y. (5). Big Data: A survey. *Mobile Networks and Applications*, 19(2), 171–209. <https://doi.org/10.1007/s11036-013-0489-0>
- [6] Katal, A., Wazid, M., & Goudar, R. H. (6). Big Data: Issues, challenges, tools and Good practices. 2013 Sixth International Conference on Contemporary Computing (IC3), 404–409. IEEE. <https://doi.org/10.1109/IC3.2013.6612229>
- [7] Sagiroglu, S., & Sinanc, D. (7). Big data: A review. 2013 International Conference on Collaboration Technologies and Systems (CTS), 42–47. IEEE. <https://doi.org/10.1109/CTS.2013.6567202>
- [8] Gandomi, A., & Haider, M. (8). Beyond the hype: Big data concepts, methods, and analytics. *International Journal of Information Management*, 35(2), 137–144. <https://doi.org/10.1016/j.ijinfomgt.2014.10.007>
- [9] Demchenko, Y., Grosso, P., De Laat, C., & Membrey, P. (9). Addressing big data issues in scientific data infrastructure. 2013 International Conference on Collaboration Technologies and Systems (CTS), 48–55. IEEE. <https://doi.org/10.1109/CTS.2013.6567203>
- [10] Hilbert, M. (10). Big data for development: A review of promises and challenges. *Development Policy Review*, 34(1), 135–174. <https://doi.org/10.1111/dpr.12142>
- [11] McAfee, A., & Brynjolfsson, E. (11). Big Data: The management revolution. *Harvard Business Review*, 90(10), 60–68.
- [12] Reinsel, D., Gantz, J., & Rydning, J. (12). The digitization of the world.
- [13] Trigyn (13) -Top Data Science and Big Data Trends to Watch in 2025.