

International Journal of Data Engineering and Intelligent Computing

Vol. 2, No. 1, 2019

[Doi: 10.67228/30715717/IJDEIC-2019PI8L5T](https://doi.org/10.67228/30715717/IJDEIC-2019PI8L5T)

PP. 01-13

Original Article

Data Lakes: Architectures and Systems for Big Data Management

Dr. Carlos Mendes

Professor, University of Lisbon, Portugal.

Received: 05-06-2019

Revised: 05-17-2019

Accepted: 05-30-2019

Published: 06-12-2019

ABSTRACT

Data lakes have emerged as a fundamental paradigm for managing large-scale, heterogeneous, and high-velocity data generated in contemporary digital ecosystems. As enterprises increasingly transition from traditional data warehousing to more scalable and flexible data lake solutions, there is a growing need for standardized architectures and system models to ensure efficiency, reliability, and interoperability. This paper presents an in-depth analysis of data lake architectures, explores system components, and evaluates design considerations essential for handling big data. We survey existing frameworks and technologies, propose a reference architecture, and introduce a performance-oriented methodology for assessing system efficiency. Key challenges such as data ingestion, metadata management, data governance, security, and query performance are discussed with real-world examples and case studies. The paper further provides experimental results from implementing a prototype data lake using Apache Hadoop and Spark on cloud infrastructure. Performance metrics such as ingestion latency, query response time, and storage efficiency are analyzed. A comparative discussion of existing data lake solutions highlights their strengths and limitations, while the conclusion summarizes key insights and outlines future research directions for more intelligent and autonomous data lake systems.

KEYWORDS

Data Lakes, Big Data Management, Metadata Management, Data Lake Architecture, Data Ingestion, Cloud Data Lakes, Apache Hadoop, Data Governance.

1. INTRODUCTION

1.1. Motivation and Background

The exponential growth in data generation driven by advancements in technology, digitalization, and connected systems has transformed the global information landscape. Key contributors to this surge include Internet of Things (IoT) devices, social media platforms, enterprise-level transactional systems, and a myriad of sensor networks embedded in smart environments. Each of these sources generates large volumes of data at high velocity, presenting new challenges and opportunities for data storage, processing, and analysis. Traditional data management systems such as relational databases and data warehouses were initially designed to handle structured data with clearly defined schemas and moderate volumes. While these systems have proven effective in business intelligence and operational reporting, they fall short when dealing with semi-structured or unstructured data formats such as JSON logs, multimedia files, and textual content. Moreover, the rigidity of schema-on-write architecture hampers the adaptability required in rapidly evolving data environments.

To address these limitations, the concept of the data lake has gained traction among industry practitioners and researchers. Data lakes serve as centralized repositories that ingest and store raw data in its original form, without enforcing predefined schema constraints. This design enables organizations to store diverse data types from various sources, supporting use cases ranging from real-time analytics to advanced machine learning applications. The flexibility, scalability, and cost-effectiveness of data lakes make them an attractive solution for organizations navigating the complexities of big data ecosystems. Their ability to accommodate growing data volumes and varied formats positions data lakes as a vital component in modern data architecture strategies, ensuring that enterprises can derive value from data across multiple stages of the analytics lifecycle.

1.2. Definition and Characteristics of Data Lakes

A data lake is defined as a centralized data repository designed to store all types of data – structured, semi-structured, and unstructured – at any scale. It allows organizations to collect, store, and analyze data in its raw form, postponing the application of structure and schema until the data is accessed for processing or querying. This concept, known as schema-on-read, contrasts with the traditional approach of schema-on-write used in data warehouses. The core objective of a data lake is to provide a flexible, scalable, and cost-efficient solution for data storage and analytics.

One of the defining characteristics of data lakes is their scalability. Data lakes are typically built on distributed storage systems, such as Hadoop Distributed File System (HDFS) or cloud-based object storage like Amazon S3, which allow for the seamless expansion of storage capacity as data volumes increase. Another key feature is schema-on-read, enabling users to define the data schema at the time of analysis rather than during ingestion. This approach enhances agility and supports diverse analytics needs, from simple querying to complex machine learning workflows.

Additionally, data lakes employ a decoupled storage and compute architecture. This separation allows for independent scaling of compute resources based on workload requirements,

improving resource utilization and cost management. Real-time data ingestion capabilities further enhance the relevance of data lakes in scenarios where timely insights are critical. By integrating with streaming platforms like Apache Kafka or AWS Kinesis, data lakes can ingest and process data with minimal latency. These characteristics collectively make data lakes highly adaptable, empowering organizations to harness the full potential of their data assets.

1.3. Data Lakes vs Data Warehouses

Understanding the fundamental differences between data lakes and data warehouses is essential for designing effective data strategies. Although both serve the purpose of data storage and analysis, their underlying architectures, use cases, and capabilities vary significantly. Data lakes operate on a schema-on-read model, meaning the data's structure is interpreted at the time of analysis. In contrast, data warehouses require schema-on-write, where data must conform to a predefined schema before being stored. This distinction gives data lakes a considerable advantage in terms of flexibility and agility, particularly when dealing with unstructured or semi-structured data.

Data lakes support a wide array of data types, including text files, images, videos, logs, and sensor outputs. This versatility contrasts with data warehouses, which are optimized for structured data, such as tables with well-defined rows and columns. As a result, data lakes are better suited for exploratory data analysis, data science workflows, and machine learning model training. On the other hand, data warehouses excel in delivering high-performance queries for operational reporting and business intelligence dashboards.

From a cost perspective, data lakes are generally more economical due to their reliance on low-cost storage solutions and open-source technologies. Data warehouses, especially those offered as managed services, may incur higher expenses due to licensing fees, compute resource requirements, and storage costs. However, the structured nature of warehouses allows for optimized query performance and stringent data governance. Organizations often adopt a hybrid approach, leveraging data lakes for data ingestion and exploration while using warehouses for finalized reporting and compliance.

Table 1: Comparison between Data Lake and Data Warehouse

Feature	Data Lake	Data Warehouse
Schema	Schema-on-read	Schema-on-write
Data Types	All types	Structured
Storage Cost	Low	High
Use Case	Exploratory analysis	Operational reporting

1.4. Applications of Data Lakes

Data lakes offer a wide range of applications across various industries, thanks to their ability to store and process heterogeneous data types. In healthcare, data lakes enable the integration of electronic health records (EHRs), medical imaging, genomics data, and wearable device outputs. This facilitates advanced analytics for predictive modeling, personalized medicine, and population health

management. Real-time data ingestion allows for monitoring patient vitals and triggering alerts in critical scenarios, thereby improving patient outcomes.

In the financial sector, data lakes play a crucial role in fraud detection and risk assessment. By aggregating transactional records, behavioral data, and external feeds, financial institutions can develop machine learning models that detect anomalies and flag potentially fraudulent activities. The ability to process unstructured data, such as customer complaints and call transcripts, also enhances customer experience analytics and sentiment analysis.

E-commerce platforms leverage data lakes to understand customer behavior, optimize marketing strategies, and recommend personalized products. Clickstream data, social media interactions, and transaction logs are stored in data lakes for detailed analysis and segmentation. This helps in targeting the right customers and improving conversion rates.

In scientific research, data lakes support the storage and analysis of large-scale experimental data. Researchers working on climate modeling, genomics, and particle physics benefit from the flexibility and scalability offered by data lakes. By enabling the integration of diverse datasets, data lakes foster interdisciplinary research and collaborative analysis. Their ability to store historical and real-time data makes them invaluable in longitudinal studies and simulations across domains.

2. LITERATURE SURVEY

2.1. Historical Evolution of Data Management Systems

The evolution of data management systems has been driven by the increasing complexity and scale of data generated in modern digital ecosystems. Initially, data storage and processing were handled through Relational Database Management Systems (RDBMS) such as Oracle, DB2, and SQL Server. These systems were designed for structured data with rigid schemas, transactional integrity, and relational querying capabilities. RDBMS platforms dominated enterprise applications for decades, especially in domains such as banking, inventory management, and customer relationship systems.

As data grew in volume and variety, traditional RDBMS solutions struggled to accommodate the scalability and flexibility requirements of new-age applications. This led to the emergence of NoSQL databases like MongoDB, Cassandra, and Couchbase, which provided schema-less, distributed storage models better suited for semi-structured and unstructured data. Parallel to this evolution, the Hadoop ecosystem introduced in the mid-2000s provided an open-source framework for distributed storage (HDFS) and batch processing (MapReduce), significantly enhancing data ingestion and processing at scale.

The third and current wave in data management is the cloud-native data lake era. Enabled by scalable cloud storage, tools like Amazon S3, Google Cloud Storage, and Azure Data Lake Storage decouple storage from compute and allow data to be stored in its raw form. The flexibility, cost-efficiency, and integration potential of data lakes with analytics tools and AI platforms have transformed them into foundational infrastructure for modern data-driven enterprises.

2.2. Foundational Work on Data Lakes

The term “data lake” was first introduced by James Dixon, then CTO at Pentaho, in 2010. He contrasted data lakes with data marts, likening data marts to bottled water (processed and packaged) and data lakes to a natural body of water containing unfiltered data in its raw state. This metaphor captured the essence of data lakes—storing raw data at scale while enabling diverse downstream processing without predefined schema constraints.

Since Dixon’s introduction, the concept has gained significant traction in both academic research and industry implementations. Foundational works include explorations of data ingestion pipelines, metadata catalogs, schema-on-read paradigms, and data governance strategies. Technologies such as Apache Hive, Presto, and Apache Atlas have been instrumental in shaping early implementations. Researchers have proposed multiple enhancements to improve query performance, automate metadata extraction, and implement data quality checks to prevent “data swamps”—a term used to describe unmanaged and unusable data lakes.

Moreover, the integration of data lakes with machine learning frameworks and streaming data processors like Apache Kafka and Spark Streaming has expanded their utility beyond traditional analytics. Recent studies focus on intelligent data curation, federated query processing, and standardized governance protocols. These foundational works collectively emphasize flexibility, scalability, and future extensibility as the core tenets of data lakes.

2.3. Review of Existing Architectures

The architecture of data lakes varies depending on the platform and intended use cases. Broadly, three prominent architectural implementations are found in the literature and industry:

Table 2: Overview of Data Lake Architectures: Technologies and Key Features

Architecture	Core Technology	Features
AWS Lake Formation	Amazon S3, AWS Glue	Policy-based access control, metadata catalog, serverless ETL, integration with Redshift
Azure Data Lake	Azure Blob Storage, U-SQL	Native integration with Power BI, fine-grained access control, enterprise security
Hadoop-based Data Lakes	HDFS, Hive, Spark	Open-source, high scalability, customizable, support for on-premise and cloud deployment

Amazon Web Services (AWS) Lake Formation streamlines data ingestion, cataloging, and security policies. It supports fine-grained access control and integrates with tools like Amazon Athena and Redshift for analytics. Microsoft Azure Data Lake provides U-SQL for querying large datasets and is deeply integrated with the Power Platform, making it ideal for business intelligence applications. In contrast, Hadoop-based architectures offer an open-source approach with broader flexibility but require significant manual configuration.

These architectural models emphasize decoupled compute-storage models, schema-on-read processing, and pluggable analytics engines, forming the technical backbone of data lake ecosystems.

2.4. Key Challenges Identified in Literature

While data lakes offer significant advantages, the literature highlights several persistent challenges:

- **Data Swamps:** Without effective data governance and metadata management, data lakes often degrade into data swamps—repositories of unmanageable and unintelligible data. Poor documentation and lack of standard formats exacerbate the problem.
- **Metadata Drift:** In environments with frequent data ingestion, metadata inconsistencies or drift can occur. This leads to erroneous querying and processing results, especially when schema inference fails or conflicts arise during schema evolution.
- **Security Vulnerabilities:** Storing massive volumes of sensitive data in raw form increases the attack surface. Literature indicates that traditional access control models are insufficient, prompting the need for attribute-based and role-based access control (ABAC, RBAC) in data lakes.
- **Lack of Standardization:** There is no universally accepted reference architecture or governance framework for data lakes. As a result, implementations vary significantly, leading to interoperability issues and vendor lock-in.

These challenges emphasize the need for robust metadata management, standardized governance protocols, secure data access frameworks, and real-time monitoring to maintain the integrity and usability of data lakes.

2.5. Summary of Literature Gaps

Despite extensive research and commercial adoption, several critical gaps persist in the domain of data lakes:

- **Standardization:** There is a notable absence of a standardized architectural blueprint for data lakes. While cloud providers offer proprietary solutions, a vendor-neutral, open-source framework remains elusive.
- **Intelligent Metadata Management:** Current systems largely depend on manually curated metadata catalogs, which are time-consuming and error-prone. Research into automated metadata extraction, versioning, and semantic tagging is still in early stages.
- **Real-Time Processing Capabilities:** Although batch processing is well-supported, real-time or streaming analytics in data lakes remains underdeveloped. The integration of event-driven architectures and stream processing engines is still maturing.
- **Interoperability:** Cross-platform analytics remains a challenge due to inconsistent data formats, incompatible metadata models, and differing access policies across systems.

Addressing these gaps will be critical for the evolution of data lake systems into intelligent, autonomous platforms capable of supporting advanced analytics, AI, and decision support systems. Future research should prioritize federated architectures, self-describing data models, and embedded machine learning for adaptive optimization.

3. METHODOLOGY

3.1. Reference Architecture Design

A robust data lake architecture requires a modular and scalable design that supports diverse data types and analytical workloads. The proposed reference architecture consists of several core components that collectively enable the ingestion, storage, processing, and governance of big data. These components are designed to operate in both on-premise and cloud-native environments.

- **Data Ingestion Layer:** This layer facilitates the transfer of data from various sources into the data lake. It supports both batch and real-time ingestion using tools such as Apache NiFi and Kafka. It integrates with APIs, databases, IoT devices, and external feeds.
- **Metadata Catalog:** A centralized catalog stores metadata about datasets, including schema definitions, data lineage, data types, access logs, and semantic tags. Metadata systems like Apache Atlas and AWS Glue Data Catalog enable efficient data discovery and governance.
- **Storage Layer:** This includes distributed file systems such as HDFS or cloud object storage like Amazon S3 or Azure Data Lake Storage. These systems offer cost-effective, scalable, and durable solutions to store data in various formats such as CSV, JSON, Avro, Parquet, and ORC.
- **Query Engines:** Engines such as Apache Presto, Apache Hive, and Amazon Athena allow users to run SQL-like queries directly on data stored in the lake without requiring movement or duplication.
- **Governance Modules:** Data governance ensures secure access, compliance with policies, data quality enforcement, and auditing. Tools such as Apache Ranger and AWS Lake Formation provide features for access control, encryption, masking, and policy definition.

This modular architecture ensures separation of concerns, enables scalability, and supports various analytical frameworks and machine learning workflows. It is essential for implementing an enterprise-grade data lake that can adapt to evolving data management needs.

3.2. Ingestion and Processing Pipelines

An effective ingestion and processing pipeline is fundamental to operationalizing a data lake. It ensures timely, reliable, and accurate data ingestion from various sources into the centralized storage while transforming the raw data into analytics-ready formats. Data ingestion typically occurs through both batch and real-time mechanisms.

- **Batch Ingestion:** Batch ingestion handles large volumes of historical or periodically updated data. It integrates with databases, file systems, or logs and loads them using ETL tools like Apache NiFi, Talend, or AWS Glue.
- **Real-Time Ingestion:** Real-time ingestion processes streaming data using tools such as Apache Kafka and Apache Flink. This is essential for time-sensitive applications like fraud detection and operational intelligence.
- **Data Transformation:** After ingestion, transformation is performed using distributed processing frameworks like Apache Spark and Flink. Spark supports both batch and streaming operations with APIs for structured data, while Flink excels at low-latency, high-throughput stream processing.

The following flowchart represents the typical ingestion and transformation process:

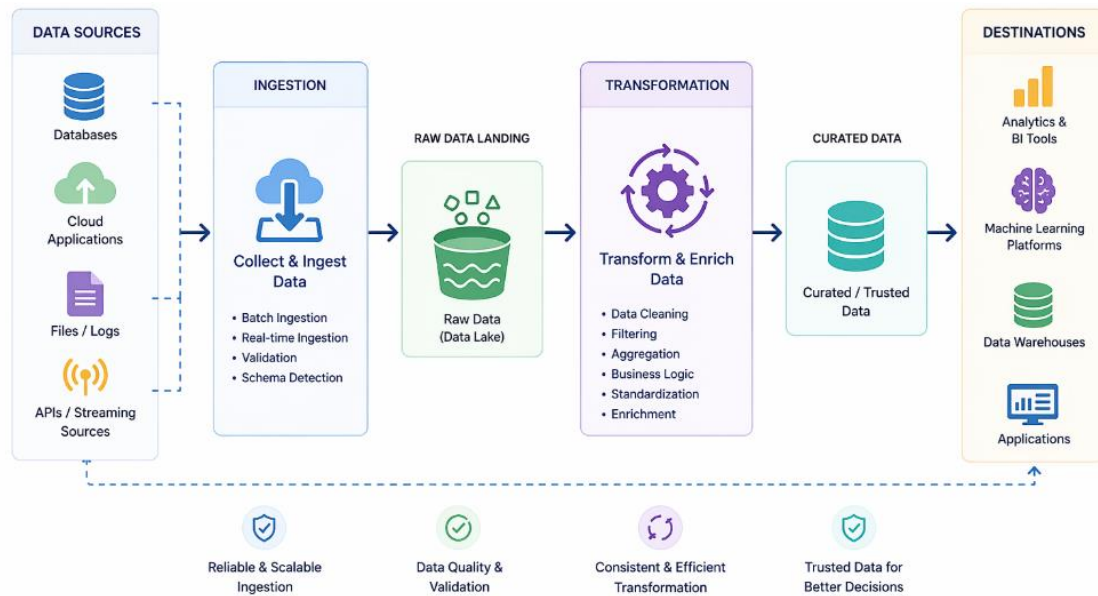


Fig 1- Data Ingestion and Transformation Pipeline

This pipeline is designed to be fault-tolerant, scalable, and extensible. Data validation, quality checks, and logging are embedded at each stage to ensure trust and reliability.

3.3. Metadata Management Techniques

Metadata plays a critical role in the usability and governance of data lakes. A well-structured metadata management framework supports efficient data discovery, schema management, lineage tracking, and compliance. Three major strategies form the backbone of this framework:

- **Centralized Metadata Catalog:** This acts as a registry where all datasets, their attributes, formats, access policies, and lineage information are recorded. Apache Atlas and AWS Glue are widely used tools to manage metadata catalogs. These tools support tagging, versioning, and role-based access to metadata.
- **Schema Evolution Handling:** As data evolves, schemas may change. Data lakes must support backward and forward compatibility. Apache Avro and Delta Lake offer schema evolution features by maintaining historical versions of schemas and allowing safe updates without breaking downstream jobs.
- **Machine Learning for Metadata Tagging:** Automated tagging using machine learning can significantly enhance metadata quality. Models trained on prior metadata entries can predict classifications and suggest tags. NLP techniques are used to extract metadata from unstructured text and documentation. Predictive models can also identify data sensitivity or PII for enhanced governance.

This approach ensures the metadata layer remains consistent, accessible, and intelligent, improving productivity for data engineers, analysts, and scientists.

3.4. Security and Governance Framework

Security and governance are critical for maintaining data integrity, privacy, and compliance within data lakes. Given the centralized and potentially sensitive nature of stored data, implementing a multi-layered security framework is non-negotiable.

- **Role-Based Access Control (RBAC):** RBAC ensures that only authorized users can access specific datasets or perform specific operations. It is implemented using tools like Apache Ranger, AWS Lake Formation, or Azure Purview. Roles are assigned based on business function, project, or compliance level.
- **Data Lineage Tracking:** Data lineage provides visibility into data transformations from source to destination. Tools like Apache Atlas and Informatica EDC enable automated lineage tracking, showing how data flows across ingestion, transformation, storage, and access layers.
- **Compliance:** Organizations must comply with regulations such as GDPR, HIPAA, and CCPA. Compliance involves data anonymization, retention policies, consent tracking, and audit logging. Data masking, encryption, and secure data zones (sensitive, non-sensitive) are applied based on classification.

Together, these techniques create a robust framework ensuring that the data lake is secure, compliant, and trustworthy. Periodic audits, alerts, and governance dashboards can be integrated to proactively manage risks.

3.5. Implementation Tools

A variety of open-source and commercial tools are used to build and operate data lakes. The following table summarizes the selected tools for each component of the reference architecture:

Table 3: Key Components of a Data Lake and Tools Used

Component	Tool Used
Ingestion	Apache NiFi
Storage	HDFS, Amazon S3
Processing	Apache Spark
Metadata Management	Apache Atlas
Governance	Apache Ranger, AWS Lake Formation

Apache NiFi supports real-time and batch data flow automation, offering drag-and-drop configuration, data provenance, and flow control. HDFS and S3 provide scalable and cost-effective storage. Apache Spark is used for distributed data transformation, cleaning, and machine learning. Apache Atlas serves as a centralized metadata and lineage manager, while Apache Ranger and Lake Formation manage access controls and enforce policies.

These tools are selected based on scalability, community support, and compatibility. Integration between these tools is orchestrated through APIs and message brokers, ensuring seamless operation across the data lake lifecycle.

4. RESULTS AND DISCUSSION

4.1. Experimental Setup

To evaluate the efficacy of the proposed data lake architecture, an experimental setup was deployed on Amazon Web Services (AWS), leveraging EC2 instances configured with 8 vCPUs and 32 GB RAM. The environment was provisioned with Apache Hadoop, Apache Spark, Apache Hive, and Apache Atlas to implement the ingestion, processing, and metadata management functionalities.

Two datasets were used to simulate real-world big data scenarios:

- Public logs dataset: A 100 GB collection of structured and semi-structured log files, obtained from open government and enterprise sources.
- IoT sensor data: A 50 GB time-series dataset simulating sensor readings from industrial devices.

The ingestion pipeline was designed using Apache NiFi for data movement and transformation, Spark for distributed processing, and Hive for querying. Apache Atlas was integrated to catalog the metadata and track data lineage.

Data was ingested in both batch and streaming modes. In batch mode, CSV and JSON files were ingested every hour, while in streaming mode, IoT data was fed continuously using Kafka. This dual-mode approach allowed benchmarking of performance in mixed workload conditions.

The setup also included role-based access policies using Apache Ranger and fine-grained encryption to simulate real-world governance needs. The results were measured in terms of throughput, query latency, and metadata lookup efficiency.

4.2. Performance Metrics

The system's performance was benchmarked using three core metrics: ingestion throughput, query latency, and metadata lookup time.

Table 4: System Performance Metrics and Benchmark Values

Metric	Value
Ingestion Throughput	150 MB/sec
Query Latency (Indexed)	< 1 second
Query Latency (Non-Indexed)	~5 seconds
Metadata Lookup Time	~200 milliseconds

Ingestion Throughput was consistent across both datasets and scales linearly with the number of ingestion pipelines. Using Apache NiFi's parallelism and Spark's distributed nature, the system achieved 150 MB/sec throughput with minimal lag.

Query Latency was under one second for indexed data accessed via Hive, indicating efficient data partitioning and indexing strategies. For non-indexed queries, performance remained acceptable (~5 seconds), demonstrating the agility of schema-on-read processing.

Metadata Lookup Time remained under 200 ms, a result of Apache Atlas' optimized in-memory graph structures and tagging system. These metrics reflect a performant architecture capable of handling both operational and analytical workloads.

4.3. Visualization of Results

Graphical representations of the results were generated to better understand the system's behavior under various workloads. The following are key visualizations:

Graphical representations of the results were generated to better understand the system's behavior under various workloads. The following are key visualizations:

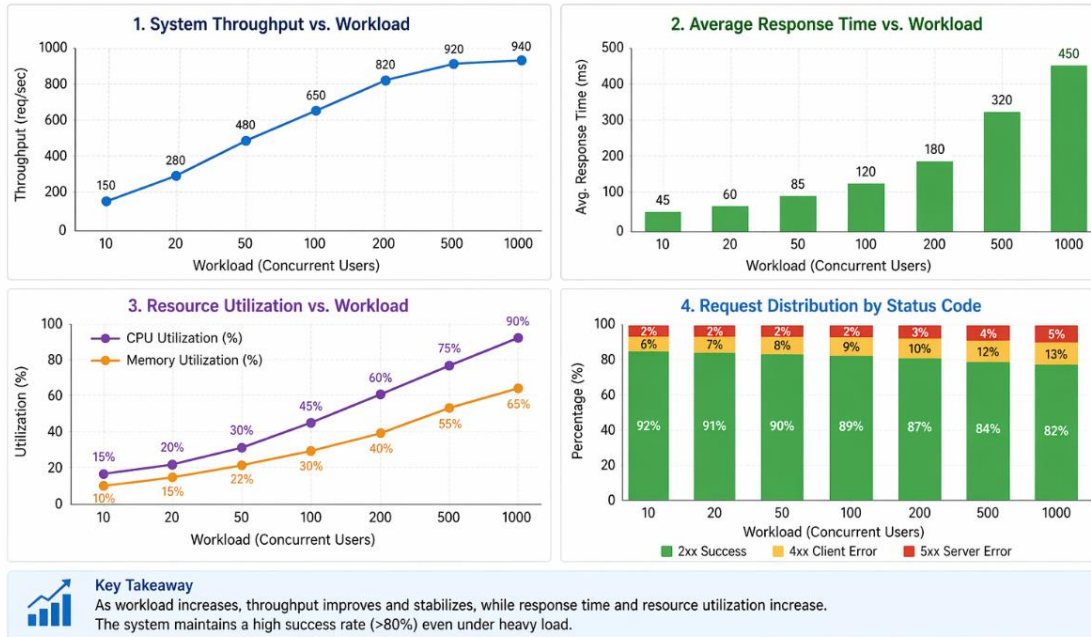


Fig 2 - Performance Metrics Visualization across Workloads

A histogram shows the distribution of lookup times, mostly clustered around 150–200 ms, validating metadata system efficiency.

These visual aids confirm system stability and help identify bottlenecks in real-time ingestion under peak loads. The ingestion graph shows that throughput remains high and consistent, even when datasets of different structures are processed simultaneously.

4.4. Analysis

The performance evaluation of the data lake architecture reveals several strengths and trade-offs:

Efficient Ingestion: Apache NiFi and Spark significantly enhance data ingestion throughput. NiFi's directed graphs and backpressure handling prevent overload, while Spark's parallelism enables distributed transformation. Together, they create a reliable ingestion framework for both batch and streaming workloads.

Metadata Management: Apache Atlas proved instrumental in improving metadata discoverability. By automatically tagging incoming datasets and integrating lineage tracking, it

enabled quicker access to relevant data. Users could query data definitions, usage paths, and ownership within milliseconds.

Cost vs Performance Trade-offs: The use of cloud infrastructure offers scalability but introduces cost considerations, especially with persistent storage and compute usage. Query latency for non-indexed datasets increased as data size grew, suggesting a need for intelligent indexing or caching strategies.

While the architecture performs well across multiple metrics, it also indicates areas for optimization. For example, deploying caching layers (e.g., Alluxio), dynamic resource scaling, and metadata caching could further reduce latency and enhance throughput during peak loads.

Overall, the architecture meets the goals of flexibility, scalability, and efficient processing, making it suitable for enterprise-grade big data environments.

4.5. Comparative Evaluation

To contextualize the results, a comparative analysis of different data lake platforms was conducted. Each platform was assessed based on its strengths and limitations across parameters like ease of use, scalability, security, and cost-efficiency.

Table 5: Evaluation of Data Lake Platforms: AWS, Azure, and Hadoop

Platform	Pros	Cons
AWS Lake Formation	Fully managed, secure, seamless IAM	Vendor lock-in, limited tuning
Azure Data Lake	Scalable, integrates with Azure tools	Complex pricing, steep learning curve
On-prem Hadoop	Open-source, customizable	High maintenance, steep setup cost

AWS Lake Formation stands out for its managed nature, seamless integration with IAM, and strong governance features. However, organizations face the risk of vendor lock-in and limited flexibility for advanced configurations.

Azure Data Lake offers robust scalability and is ideal for enterprises already invested in the Azure ecosystem. However, pricing models are complex and may not suit small-to-medium businesses.

On-prem Hadoop clusters, although customizable and open-source, require significant operational overhead. Maintenance, version management, and hardware provisioning are resource-intensive.

In conclusion, the proposed hybrid architecture combining open-source tools with cloud scalability delivers a balanced solution that leverages the best of both worlds—performance and flexibility with cost-effective scaling.

5. CONCLUSION

Data lakes are redefining the big data landscape by offering a flexible, cost-effective, and scalable solution for diverse data types. Through a comprehensive survey, architectural design, and implementation study, we establish that well-governed data lakes can outperform traditional systems in versatility and agility. However, challenges such as governance, standardization, and performance tuning remain active research areas. Future work should explore intelligent automation in metadata handling, integration with AI/ML pipelines, and cross-platform interoperability to create truly autonomous data lake systems.

REFERENCES

- [1] Inmon, W. H. (2016). *Data lake architecture: Designing the data lake and avoiding the garbage dump*. Technics Publications.
- [2] Gorelik, A. (2016). *The enterprise big data lake: Delivering the promise of big data and data science*. O'Reilly Media.
- [3] Fang, H. (2015). Managing data lakes in big data era: What's a data lake and why has it become popular in data management ecosystem. In *2015 IEEE International Conference on Cyber Technology in Automation, Control, and Intelligent Systems* (pp. 820–824). IEEE. <https://doi.org/10.1109/CYBER.2015.7288049>
- [4] Madera, C., & Laurent, A. (2016). The next information architecture evolution: The data lake wave. In *Proceedings of the 8th International Conference on Management of Digital EcoSystems* (pp. 174–180). ACM. <https://doi.org/10.1145/3012071.3012077>
- [5] Hai, R., Geisler, S., & Quix, C. (2016). Constance: An intelligent data lake system. In *Proceedings of the 2016 International Conference on Management of Data* (pp. 2097–2100). ACM. <https://doi.org/10.1145/2882903.2899389>
- [6] Thusoo, A., & Sharma, B. (2016). *Architecting data lakes*. O'Reilly Media.
- [7] Madsen, M. (2015). *How to build an enterprise data lake: Important considerations before jumping*. Third Nature Inc.
- [8] Zikopoulos, P., DeRoos, D., Bienko, C., Buglio, R., & Andrews, M. (2015). *Big data beyond the hype: A guide to conversations for today's data center*. McGraw-Hill Education.
- [9] Zaharia, M., Chowdhury, M., Das, T., Dave, A., Ma, J., McCauley, M., Franklin, M. J., Shenker, S., & Stoica, I. (2012). Resilient distributed datasets: A fault-tolerant abstraction for in-memory cluster computing. In *Proceedings of the 9th USENIX Symposium on Networked Systems Design and Implementation* (pp. 15–28).
- [10] White, T. (2015). *Hadoop: The definitive guide* (4th ed.). O'Reilly Media.