

Data Lake Architecture for Smart Fish Farming Data-driven Strategy

Mohamed El Mehdi El Aissi, Younes Lakhrissi, Safae El Haj Ben Ali

University of Sidi Mohamed Ben Abdellah - Siger Laboratory

sarah.benjelloun@usmba.ac.ma

Abstract. Thanks to continuously evolving data management solutions, data-driven strategies are considered the main success factor in many domains. These strategies consider data as the backbone, allowing advanced data analytics. However, in the agricultural field generally, and especially fish farming, data-driven strategies have yet to be widely adopted. This research paper aims to demystify the situation of the fish farming domain in general by shedding light on big data generated in fish farms. The purpose is to propose a dedicated data lake functional architecture and extend it to a technical architecture to initiate a fish farming data-driven strategy. The research opted for an exploratory study to explore the existing big data technologies and propose an architecture applicable to the fish farming data-driven strategy. The paper provides a view of how big data technologies offer multiple advantages for decision-making and enabling prediction use cases. Also, it highlights different Big Data Technologies and their use. Finally, the paper presents the proposed architecture to initiate a data-driven strategy in the fish farming domain.

Keywords: Data-driven Strategy, Big Data, Data Lake, Fish Farming, Data Analytics

1. Introduction

Nowadays, Big Data and its applications are ubiquitous since businesses and organizations have increasingly adopted data-driven strategies which allow them to tackle many challenges (Sawant, 2013). In parallel, cloud computing and high-performance computing became increasingly accessible as a service, which made adopting Big Data solutions in many domains possible (Nachiappan, 2017). However, despite all the advantages of adopting a data-driven strategy, the agricultural domain did not benefit sufficiently from it yet, especially the fish farming fields (Maru, 2018). Moreover, compared to other domains, the agricultural field, in general, and fish farming, in particular, has a high level of uncertainty due to various parameters that affect the quality and quantity of production (Li, 2011). Consequently, applying a data-driven strategy in the fish farming domain will not only bring benefits to farmers but also enhance the entire sector as a whole. The proposition to adopt a strategy for the whole domain instead of focusing on the fish farms specifically has not been discussed yet and the proposition of an architecture to initiate this strategy has not been addressed either. This assertion is a result of a literature and comprehensive review more detailed in the methodology section.

A data-driven strategy for fish farming will reduce uncertainties and provide continuous decision support, revolutionizing the fish farming sector (Maru, 2018). Nevertheless, making the data the pillar of this strategy will directly impact the production efficiency and nutrition quality with the minimum of environmental damages, which implies social, economic, and ecological benefits (Elgendy, 2014).

To obtain crucial information about massive datasets, advanced data analysis techniques have been used to transform big data into valuable data since it helps organizations and businesses make better decisions. For example, in the field of healthcare analytics, large datasets (provided by applications such as Electronic Health Records and Clinical Decision Systems) may enable healthcare practitioners to provide effective and accurate solutions for patients based on their overall history rather than relying only on collected current data (Pramanik, 2022). It has to be noted that traditional data analytics are deprecated to use in a big data context as long as the V's that characterize big data, namely, veracity, value, velocity, volume and variety, will immediately affect the accuracy of results. Additionally, other characteristics may be added to big data as viscosity, venue, validity, and viability (Panimalar, 2017). Furthermore, many artificial intelligence (AI) methods were developed to extract valuable information and discover hidden patterns from enormous amounts of data more accurately and faster than traditional techniques (Mahesh, 2020), such as machine learning (ML) or data mining (DM). For instance, a detailed analysis of historical fish farming data, real-time collected data and other parameters will provide a detailed and extensive vision of the fish farming market status, thereby taking the appropriate decisions (Wang, 2021).

This article is organized into seven sections—first, an introduction to present big data applications and how they can be used to enhance domains. Second, we present the methodology used in this study. After that, we highlight in section three the advantages of adopting data-driven strategies in several fields and how it enables the adoption of fish farming data-driven strategies. The fourth section presents the different big data technologies used for proposing a dedicated data lake architecture. Section five proposes a technical data lake architecture for handling fish farming data based on the Hadoop ecosystem. Section six is a discussion to shed light on the positive impact of adopting a data-driven strategy based on data lake architecture. Finally, section seven resumes the different axes tackled in this paper then it describes future work that includes developing a proof of concept based on the proposed fish farming data lake architecture.

2. Methodology

The primary research method for this study is literature review. The objective is to identify sources that address the same research question : How big data technologies can help enhance the fish farming domain ? All literature available and published after 2018 have been assessed. Other than the period criteria, we used a few criteria to narrow down the sources: scholarly (peer reviewed) sources; full article publication; language publication in English; articles focusing on technical design. To collect articles from renewed research databases such as web of science, Springer and IEEE Xplore, we used the following query: ["Big Data" OR "Data-driven" OR "Big Data Technologies" OR "Data Lake Architecture"] AND ["Aquaculture" OR "Fish Farming"]. It is important to note that this research is limited by manpower and time thus; the sources found may not be exhaustive.

Following, we examined the academic sources based on a comprehensive review in order to identify the current industry data-based solutions. The remaining sources address the following subjects:

- Review of Big Data in aquaculture.
- A data science method proposal to address water management, disease detection, feeding strategies, fish behavior monitoring.

These sources focus on problem resolution on the farm level. The proposition of architecture to initiate a data-driven strategy for the fish farming domain is yet to be addressed.

In the next step of this study, we followed an exploratory study to explore the existing big data technologies and propose an architecture applicable to the fish farming data-driven strategy.

3. Big Data Enabling the Adoption of Fish Farming Data-Driven Strategies

The amount of available data does not solely determine the big data value but is determined by how it has been used. Analyzing data from different sources increases operational efficiencies, upgrading product development and driving new revenue and growth prospects by enabling smart decision-making.

Big Data technologies and their applications have been adopted in many fields and have earned one's spurs. Indeed, it offers multiple advantages in terms of detecting the leading causes of failures/anomalies, evaluating the potential risks that may occur, spotting fraudulent situations before it harms the business or increasing the accuracy of artificial intelligence models. Below is a detailed review shading lights on the multiple advantages that big data offers by domain (Sagiroglu, 2013):

- **Public sector:** increasing transparency through accessible, connected data, identifying demands, customizing actions for appropriate products and services especially, decision making with fewer risks and inventing modern services and products (Coulthart, 2022).
- **Industry:** precise demand forecasting, higher supply chain planning, sales support, advanced production operations and integrating web search-based applications (Li et al., 2022; Yoon et al.2020).
- **Marketing:** near real-time clients behavior analysis based on marketplace collected data, geo-targeted advertising, more effective product design, price and variety optimization, distribution and logistics management (Cao, 2022).
- **Healthcare:** clinical decision support systems, personal analytics based on each patient profile history, personalized treatment, illness pattern analysis, and performance-based payment for medical staff (Rehman et al. 2022; Hussein et al. 202).

Making data at the center of a strategy will provide immense opportunities for organizations to improve (Lusch et al., 2015; Rajaraman, 2016). Adopting the same concept in the agricultural domain and especially in the fish farming sector will enhance the domain as a whole by handling multiple parameters that could tackle the development of the fish farming domain (Mouzakitis et al., 2020). However, designing a fish farming data-driven strategy can only be possible by implementing a dedicated big data architecture, offering the ability to handle heterogeneous data from different sources.

Handling Fish farming data is not easy as it appears since the objective is to implement a data-driven strategy whose cornerstone is built upon data. Nevertheless, the collected and generated data records have the characteristics of Big Data (the V's of big data), so handling it is challenging. As a consequence, designing a dedicated Data Lake architecture is primordial.

The Data Lake remains an adequate data handling platform for the fish farming use case, as it relies on the Extract Load Transform (ELT) process rather than the Extract Transform Load (ETL) process adopted in data warehouse architectures (Nambiar, 2022). The ELT process allows ingesting data from heterogeneous sources without undergoing a specific schema and then processing the data to extract hidden patterns and valuable information before exposing it to the end user (Aissi, 2022).

The Data Lake architecture has evolved from mono-zone architecture to multi-zone architecture (Ravat, 2019). With the mono zone concept, the data lake has a flat architecture where all the collected data is stored in their raw format, with low costs. Despite that, the mono zone data lake architecture has limited user access and restricted data processing operations. The multi-zone concept usually relies on three main zones, namely, the RZ (Raw Zone), CZ (Conformed Zone) and AZ (Analytics Zone). The RZ is where the data is stored as-is without any transformation; the ingestion can be in batch mode or real-time mode. The main goal of this zone is to provide data engineers with the original version of data and facilitate subsequent processing. The previously stored data is transformed at the CZ level to match specific needs. Also, it has to be noted that both batch and real-time processing can be adopted to make data accessible for analysis purposes. Processing includes operations such as selecting, joining, and aggregating datasets (Benjelloun, 2020). The AZ, also known as the access zone, exposes the preprocessed data and access to it is granted to data analysts. This zone provides a self-service data exploration for the analytics use cases (business intelligence, machine learning algorithms, statistical analysis, and reporting).

Along with this, our proposed functional architecture for handling fish farming big data contains three main phases, namely, data generation & collection, data ingestion & treatment and data analysis & exploration, as presented in the figure below:

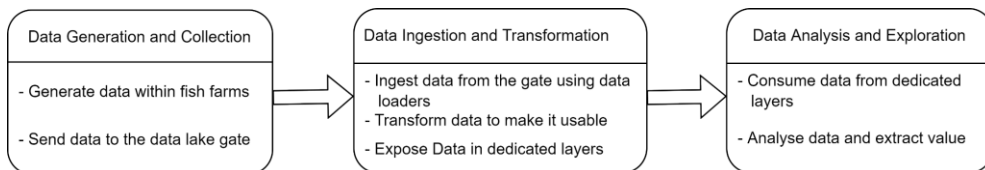


Fig. 1: Process of Handling Fish Farming Data

By concretizing this functional data handling architecture, the fish farming domain will not remain a traditional field but a data-driven domain, thanks to the possibility of analyzing and extracting valuable information from massive data gathered from multiple fish farms (Sarah et al., 2022).

In other terms, the fish farming data lake will make data accessible for multiple users to perform analysis and not restrict access to only users with a solid technical

background. Adopting a data-driven strategy will become possible, thus enhancing the whole domain.

4. Big Data Technologies for Fish Farming Use Case

4.1. Hadoop Ecosystem

Apache Hadoop is a popular big data technology with a large community. The main goal of this technology is to tackle the various challenges of performance and complexity while working with massive data using traditional systems. Hadoop is designed to effectively perform advanced processing on colossal data sets through a distributed file system (Vohra, 2016). Indeed, Hadoop runs the data processing tasks where the data are stored in the cluster and not copied in memory. As a result, Hadoop allows querying terabytes of data in a few seconds, with high fault tolerance, as it replicates data on servers to avert data loss (Monteith, 2018).

Hadoop relies on two main components to provide a robust big data platform: Hadoop Distributed File System (HDFS) and the MapReduce (MapR) model. Moreover, depending on the need, we may install new modules on top of Apache Hadoop to match applications' requirements.

4.2. Hadoop Distributed File System (HDFS)

Hadoop Distributed File System is based on master-slave architecture to store massive data files with high scalability and low cost. It can store structured, unstructured and semi-structured data. The main advantage of HDFS is the concept of delegating computation tasks near to data location since it reduces the network congestion in the cluster. Indeed, the cluster is constituted of two types of servers, the Name Node (master) that is responsible for managing file systems operations and multiple Data Nodes (slaves) that coordinate data storage and computations (Oussous, 2018).

4.3. MapReduce Model (MapR)

MapReduce is a programming model implemented on top of HDFS. It is the cornerstone of big data management as it allows the efficient processing of huge amounts of data. The MapR model relies on parallel processing and contains two main functions: the Map and the Reduce (Condie, 2010).

In the Map function, there are two actions, first splitting the dataset into equal units or chunks constituting key-value pairs. The key-value pairs are given to the mapper that runs parallel mapping tasks across the cluster and resulting intermediate key-value pairs. It has to be noted that the resulting key values are sorted and grouped before transmitting them to the Reduce phase.

In the Reduce function, process the intermediate key-value pairs. For each key, the Reduce function aggregates the values with a coding logic to provide a summary of the entire dataset. Finally, the output is stored in the HDFS.

4.4. Data Exposition: Hive, Hbase, Elasticsearch

Apache Hive is a distributed, fault-tolerant data warehouse system built on top of HDFS. It is designed to efficiently store and analyze huge datasets by centralizing them in a single store. Indeed, Hive is based on structured tables. Each table has a related HDFS directory that is divided into partitions, and each partition is divided into buckets. In addition, for better data analysis, Hive provides a SQL-Like language named HiveQL for querying data. Each HiveQL query is converted into MapR jobs processed in parallel and impacting data in HDFS directly (Shaw, 2016).

Apache HBase is a distributed, column-oriented, non-relational database with a key/value model built on top of HDFS. It is designed to support high table-update. In addition, HBase allows gathering multiple data elements into column families with a unique row key. However, HBase is most effectively used to store non-relational data in data-driven strategies. The HBase API enables random, strictly consistent, real-time access to large volumes of data. In fact, HBase tables are known as HStore, and each HStore has one or more associated Map-Files stored in HDFS (Prasad, 2016).

Elasticsearch is a NoSQL, document-oriented database. It stores data in an unstructured way and cannot be queried using SQL. Since it is based on indices, the entire object graph needs to be indexed to be searched. Elasticsearch is a distributed search and analytics engine. Kibana is proprietary data visualization dashboard software for Elasticsearch. Kibana enables to interactively explore data in Elasticsearch using Kibana Query Language to filter data using free text or field-based search. The elastic stack can be used not only for functional data but also to Store and analyze logs, metrics, and security event data (Elasticsearch, 2019).

4.5. Data Ingestion: Apache Sqoop, Apache Flume

Apache Sqoop is a tool designed for efficiently transferring bulk data between Apache Hadoop and relational database management systems such as MySQL and Oracle. It is an open-source command-line interface application ETL tool (Lakhe, 2016).

Apache Flume is an open-source, powerful and flexible system to collect, aggregate, and move large amounts of unstructured data from multiple data sources into Hadoop (HDFS or HBase). It is written in Java and has a flexible architecture based on streaming data flows. Flume has its own query processing engine to transform new batches before moving to the intended sink (Lakhe, 2016).

4.6. Data Processing: Apache Spark

Apache Spark is an open-source distributed data processing engine. Unlike the MapR framework, Spark uses in-memory caching to optimize performance. In addition, Spark enables complex processing on huge datasets through development APIs in Scala, Python or Java. Indeed, Apache Spark relies on the Resilient

Distributed Dataset (RDD) concept, which is a constructed collection of data from a source system or another RDD stored in-memory. In case of failure, Spark can reconstruct the RDD thanks to a Direct Acyclic Graph that describes the sequence of operations, hence the resilience of RDDs. It has to be noted that while working on DataFrames and Datasets, Spark implicitly transforms them to RDD before performing any processing. The Spark framework includes many components, such as those listed below (Grishchenko):

- **Spark Core:** It is the cornerstone of the platform as it is responsible for managing memory, distributing tasks, interacting with HDFS and fault recovery.
- **Spark SQL:** It is a distributed engine allowing fast, interactive querying compared to the MapR model. Moreover, it can write the output RDD to Hive tables and use HiveQL.
- **Spark Streaming:** It is a real-time engine allowing streaming data analysis. Mainly, it relies on the micro-batch data ingestion concept and enables data processing with almost the same logic as batch processing.

4.7. Stream Processing: Apache Kafka

Apache Kafka is an open-source distributed event store and stream-processing platform written in Java and Scala. Kafka provides a low latency, high throughput platform for handling real-time data feeds. It guarantees zero data loss and is very fast, as it performs over 2 million writes per second. Kafka is used to building real-time streaming pipelines to move data from one system to another. Indeed, Kafka is a publish-subscribe messaging system that receives data from multiple sources and makes it available for listening systems. An application publishes a stream of events to a topic on a Kafka broker, which can then be consumed independently by other applications (Bandi, 2021).

4.8. Data Scheduling: Apache Oozie, Apache Airflow

Oozie is a server-based workflow scheduling system to manage Hadoop jobs. The workflows on Oozie are based on XML (extensible markup language) and defined using control flow and action nodes in a Directed Acyclic Graph (DAG) (Islam et al., 2015).

Airflow is an open-source workflow management platform for data engineering pipelines developed at Airbnb to manage workflows' increasing complexity. It is based on Python, and it authors workflows as DAGs of tasks.

5. Technical Data Lake Architecture Proposal for Fish Farming Data-Driven Strategy

With the emergence of data-driven strategies in many fields and the massive increase in the quantity of generated data by the fish farming domain, it is

imperative to design data lake architecture able to handle data coming from different sources and offer a solid platform to perform advanced analytics. The proposed fish farming data lake architecture, represented in Figure 2, represents a flexible, scalable and efficient solution that covers the entire process of collecting, transforming and exposing data.

To ease the implementation of this architecture and make it widely accessible, the fish farming Data Lake was designed according to six main guidelines:

1. It should support handling the three types of data (structured, semi-structured, unstructured)
2. The process of collecting and ingesting data should be automated
3. Depending on the provenance and type, the data should be classified
4. It should support handling historical/ backup data in raw format
5. Providing a dedicated data access layer for data analyst teams
6. Providing a data science laboratory for advanced analytics

In the same perspective, the proposed fish farming data lake architecture is multi-zone architecture with three main zones; namely, Raw Zone (RZ), Trusted Zone (TZ) and Access Zone (AZ).

The data could be gathered from different sources such as Relational Data Base Management Systems (RDBMS), CRM/ERP, Human Generated Data such as flat files, Application Programming Interfaces (APIs) like weather and IoT/Sensors for data like Ph, temperature and pressure. Then, directly submitted to the data lake edge node, considered the entry point.

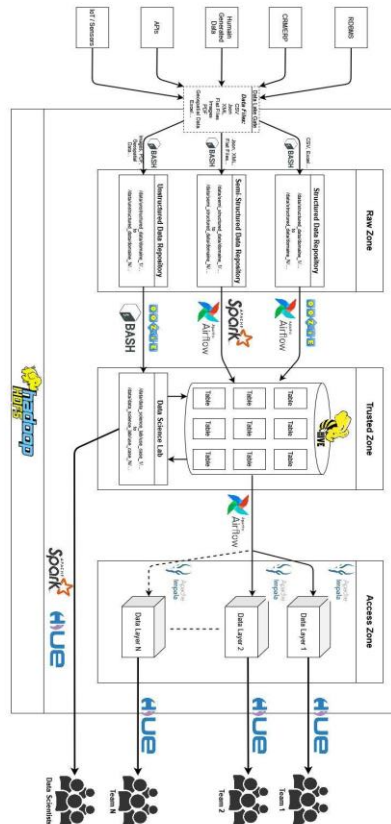


Fig. 2: Fish Farming Data Lake Technical Architecture

5.1. Raw Zone

The gateway represents the entry point of the data lake. The main objective of the gateway machine is extracting data from multiple sources before loading it into the Hadoop Distributed File System. Each data source has access to this server. However, this server should not be confused with HDFS. The received data files follow a nomination pattern that indicates their source and date. The reason behind using this approach is to allow keeping track of the collected data.

The data loader is the mechanism that allows the loading of data to HDFS. It is a developed job that reads the previously received data files from the gateway and stores them in the data lake file system. It relies on a defined naming pattern to identify the source, domain and time of data batches then it loads each one in the appropriate HDFS repository. In addition, the data loader can distinguish between structured, semi-structured and unstructured data based on the data file extension. On the HDFS side, three repositories are created: structured, semi-structured and unstructured data inside each directory, data is organized by source, and we have different domains inside each source. These three HDFS repositories constitute the Raw Zone. It must be noted that the data is stored as-is from the source at this level.

5.2. Trusted Zone

To pass this data to the trusted zone, we rely on different jobs depending on the data type. For structured data, it is directly stored in structured tables. Apache hive is a data warehouse system based on two types of tables: External and managed. External tables point directly to a remote data directory with a defined structure (fields, delimiter, etc.). Moreover, it has to be noted that Hive does not manage the storage of the external table; it only manages the metadata, which means, in case of deleting the external table, the data remains in HDFS, and only the table definition is deleted. For managed tables, also known as internal tables, both the storage and the metadata are managed by Hive, and deleting the table implies deleting the data from HDFS. The table creation is a one-time action. Still, the data ingestion into the conformed zone is an automated process using an orchestrator such as Apache Oozie that allows the execution of sequential and parallel tasks with a defined frequency to match the frequency of data collection.

Semi-structured data is not collected in a conventional or tabular form but is not completely unstructured. It contains some tags or key values that can be used to analyze that data. In our data lake use case, semi-structured data may refer to JSON files, XML files or other markup language-based files. This data is transformed using dedicated Apache Spark jobs to match a defined structure at the level of the conformed zone. A dedicated spark job is developed for each semi-structured data source to transform and store data values in the conformed zone in a structured hive table. The spark jobs are executed systematically using the orchestrator to ensure a continuous semi-structured data integration.

In the fish farming data lake use case, unstructured data, such as images and geospatial data, is usually used for data science use cases. For this purpose, data transition scripts are developed to create dedicated data science repositories for each use case.

5.3. Access Zone

The access zone objective is to allow easy access to all data available on the data lake for reporting, dashboarding or data analysis purposes. In this optic, a dedicated layer is created for each need with controlled data access. These layers are either a group of hive tables or specific directories containing unstructured data. The data access is controlled using Apache Ranger policies. By this architecture, we can ensure that all external tools/teams can plug into the data lake and read the specified data.

6. Discussion

The agricultural sector produces massive data that can enhance the sector as a whole and anticipate market needs. In addition, through big data technologies, advanced data analytics platforms are designed to facilitate the process of extracting valuable

information from stored data for farmers and researchers (R. Lokers, 2016). For instance, weather forecasting can be done based on the collected data from APIs since it is considered a key factor for fish farming. Moreover, it is very interesting for fish farmers to continuously measure tank temperature and Ph, which can be easily retrieved from data. In addition, big data technologies offer huge opportunities for farmers to adopt better fish management by analyzing different stages of fish farming, such as fingerling, nutrition and production time (M. R. Bendre, 2015).

Another critical point is to analyze the market information to get profit from fish production. It has to be noted that many parameters can constitute market data, such as input cost, price trends, farming cost, demand and supply, and transportation cost. Furthermore, this data may be accessible to other public and private parties (Sarker, 2020).

Given these points, the proposed fish farming data lake architecture is designed to initiate the adoption of a data-driven strategy based on the results of analyzing huge data existing on the data lake. This architecture is designed to collect data existing in different farms with different formats and types. Then, it is directly saved into the RZ without applying any transformation to it, with the aim of forming a solid historical repository of the received data. Then comes the TZ, where lean transformations are applied to make data ready for further analysis. Finally, the AZ is made up of a dedicated layer with the required data for each team.

7. Conclusion

In this paper, we discuss the positive impact that data-driven strategies have when adopted in a domain. In addition, we take several domains, such as marketing, healthcare and the industry of construction, as a reference to confirm that making data the center of a strategy allows extracting valuable information to support decision-making and predictive analysis. However, we shed light on the agricultural domain as it does not benefit sufficiently from the advantage of adopting a data-driven strategy, especially in the fish farming field.

In addition, designing dedicated data lake architecture is the cornerstone of adopting a fish farming data-driven strategy. Furthermore, we propose a complete fish farming big data architecture for collecting, handling and processing data, based on a multi-zones Data Lake architecture (RZ, CZ, AZ), allowing huge data analysis with high efficiency and scalability. Moreover, we present the previously proposed data lake functional architecture from a technical perspective by demystifying each Hadoop ecosystem component used while implementing this solution. On top of that, we explain the technical use of each one of the mentioned data lake zones.

Ultimately, our future work will focus on implementing this proposed dedicated

fish farming data lake architecture to provide a proof of concept (POC) and simulate the data collection, transformation and exposition process based on raw generated fish farming data.

Acknowledgement

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Conflict of Interest Statement

The authors have no conflicts of interest to declare. There is no financial interest to report, and all co-authors have seen and agree with the manuscript's contents.

References

- Sawant, Nitin, & Himanshu Shah. (2013). Big Data Application Architecture. Big data Application Architecture Q & A. Apress, Berkeley, CA. 9-28
- Nachiappan, R. et al. (2017). Cloud storage reliability for big data applications: A state of the art survey. *Journal of Network and Computer Applications*, 97: 35-47.
- Maru, A. et al. (2018). Digital and data-driven agriculture: Harnessing the power of data for smallholders. *F1000Research* 7.525, 525
- Li, Xuepeng, et al. (2011). Aquaculture industry in China: Current state, challenges, and outlook. *Reviews in Fisheries Science*, 19.3, 187-200
- Pramanik, P., Kanti, D., Saurabh, P., & Moutan, M. (2022). Healthcare big data: A comprehensive overview. *Research Anthology on Big Data Analytics, Architectures, and Applications*, 119-147
- Panimalar, A., Varnekha, S., & Veneshia, K. (2017). The 17 V's of big data. *International Research Journal of Engineering and Technology (IRJET)* 4.9, 3-6
- Mahesh, B. (2020). Machine learning algorithms-A review. *International Journal of Science and Research (IJSR)*. [Internet] 9, 381-386
- Sagiroglu, S. & Duygu, S. (2013). Big data: A review. 2013 International Conference on Collaboration Technologies and Systems (CTS). IEEE
- Aissi, El, et al. (2022). Data Lake Versus Data Warehouse Architecture: A Comparative Study. WITS 2020. Springer, Singapore, 201-210
- Ravat, F. & Yan, Z. (2019). Data lakes: Trends and perspectives. *International Conference on Database and Expert Systems Applications*. Springer, Cham

Benjelloun, S. et al. (2020). Big data processing: Batch-based processing and stream-based processing. 2020 Fourth International Conference on Intelligent Computing in Data Sciences (ICDS). IEEE

Monteith, J. Yates, John D. McGregor, and John E. Ingram. (2013). Hadoop and its Evolving Ecosystem. 5th International Workshop on Software Ecosystems (IWSECO 2013), 50

Oussous, A., et al. (2018). Big Data technologies: A survey. *Journal of King Saud University-Computer and Information Sciences* 30(4), 431-448

Condie, T., et al. (2010). MapReduce online. *Nsdi*, 10(4)

Shaw, S., et al. (2016). *Hive Architecture. Practical Hive*. Apress, Berkeley, CA, 37-48

Prasad, B. R. & Sonali, A. (2016). Comparative study of big data computing and storage tools: A review. *International Journal of Database Theory and Application* 9(1), 45-66

Elasticsearch, B. V. (2018). Elasticsearch. Internet: <https://www.elastic.co/pt/>, [Sep. 12, 2019]

Lakhe, B. (2016). Implementing SQOOP and Flume-based Data Transfers. *Practical Hadoop Migration*. Apress, Berkeley, CA, 189-205

Grishchenko, A. *Spark Architecture*

Bandi, A., & Hurtado, J. A. (2021). Big data streaming architecture for edge computing using kafka and rockset. 2021 5th International Conference on Computing Methodologies and Communication (ICCMC). IEEE

Lokers, R., Knapen, R., Janssen, S., van Randen, Y., & Jansen, J. (2016). Analysis of Big Data technologies for use in agro-environmental science. *Environ. Model. Softw.*, 84, 494–504

Bendre, M. R., Thool, R. C., & Thool, V. R. (2015). Big data in precision agriculture: Weather forecasting for future farming. In 2015 1st International Conference on Next Generation Computing Technologies (NGCT), 744–750

Sarker, M. N. I., Islam, M. S., Murmu, H., & Rozario, E. (2020). Role of big data on digital farming. *Int J Sci Technol Res*, 9(4), 1222-1225

Elgendy, N., & Elragal, A. (2014, July). Big data analytics: a literature review paper. In *Industrial conference on data mining* (pp. 214-227). Springer, cham

Wang, C., Li, Z., Wang, T., Xu, X., Zhang, X., & Li, D. (2021). Intelligent fish farm—the future of aquaculture. *Aquaculture International*, 29(6), 2681-2711

Rehman, A, Saeeda, N, & Razzak, A. (2022). Leveraging big data analytics in healthcare enhancement: trends, challenges and opportunities. *Multimedia Systems* 28(4), 1339-1371

Li, C., Chen, Y., & Shang, Y. (2022). A review of industrial big data for decision making in intelligent manufacturing. *Engineering Science and Technology, an International Journal*, 29, 101021

Cao, G., Tian, N., & Blankson, C. (2022). Big data, marketing analytics, and firm marketing capabilities. *Journal of Computer Information Systems*, 62(3), 442-451

Coulthart, S., & Riccucci, R. Putting big data to work in government: The case of the United States border patrol. *Public Administration Review*, 82(2), 280-289

Lusch, R. F. & Nambisan, S. (2015). Service innovation. *MIS quarterly*, 39(1), 155-176

Rajaraman, V. (2016). Big data analytics. *Resonance*, 21(8), 695-716

Nambiar, A. & Mundra, D. (2022). An Overview of Data Warehouse and Data Lake in Modern Enterprise Data Management. *Big Data and Cognitive Computing*, 6(4), 132

Mouzakitis, S., et al (2020). Investigation of common big data analytics and decision-making requirements across diverse precision agriculture and livestock farming use cases. *International symposium on environmental software systems*. Springer, Cham

Vohra, D. (2016). *Practical Hadoop ecosystem: A definitive guide to hadoop-related frameworks and tools*. Apress

Islam, M. K. & Aravind, S. (2015). *Apache Oozie: The Workflow Scheduler for Hadoop*. O'Reilly Media, Inc

Benjelloun, S., Aissi, M. E. M. E., Lakhri, Y., Ali, S. E. H. B. (2022). Big Data Technology Architecture Proposal for Smart Agriculture for Moroccan Fish Farming. *WSEAS Transactions on Information Science and Applications*, 19, 311-322

Yoon, J. & Joung, S. (2020). A big data based cosmetic recommendation algorithm. *Journal of System and Management Sciences*, 10(2), 40-52

Hussein, A., Ahmad, F. K., & Kamaruddin, S. S. (2021). Cluster Analysis on covid-19 Outbreak Sentiments from Twitter Data using K-means Algorithm. *Journal of System and Management Sciences*, 11(4), 167-189