

**APACHE SPARK-BASED DISTRIBUTED FRAMEWORK FOR SCALABLE EDI
TRANSACTION PROCESSING.**

¹.Naga Sai Mrunal Vuppala, ².Parth joshi , ³. Harsh Parnerkar, ⁴Shalmali Joshi

¹Senior Software Engineer

Independent Researcher, Dallas, TX, USA

ORCID: <https://orcid.org/0009-0007-6389-6544>

Email: mrunalvppl@gmail.com

²Designation: Software Engineer

Affiliation: Tata Consultancy Services

Country:USA

Gmail ID: parthuta@gmail.com

<https://orcid.org/>

0009-0001-7000-3501

³Designation: Software Engineer II

Affiliation: JP Morgan Chase

Gmail: hp.parkerkar@gmail.com

Orchid: 0009-0003-4317-924X

⁴Advanced Analytics Consultant

Affiliation: Elevance Health

Gmail: joshishalmali@gmail.com

ORCID: 009-0000-4329-3841

Abstract

The article describes how to use Apache Spark to improve Electronic Data Interchange (EDI) transactions. Apache Spark has high potential to address scalability, performance, and integration issues that commonly occur during peak loads in traditional EDI systems, including retail and logistics settings. In this study, it was established that the distributed computing services offered by Spark minimise load, increase processing speed, and minimise resource utilisation. The research compares performance by developing a quantitative assessment of real-world data from EDI systems using Spark and traditional centralised systems. Analysis results indicate that Spark-based systems are quick at handling transaction processing, can be extended, and are cost-effective. The article is about modernising the business EDI infrastructure, and as a result, massive savings have been made. It also implies directions for future research to make the use of distributed systems even lighter in EDI applications, e.g., Apache Spark.

Key Words: *Apache Spark, Electronic Data Interchange (EDI), Distributed Computing, Scalability, Transaction Processing.*

1. Introduction

Organizational-to-organizational exchange of business documents via EDI is in use, e.g., invoices, purchase orders, or shipping notices. EDI is a critical component of business today, especially in manufacturing and logistics. Statistics indicate that a large percentage of organizations worldwide use EDI to conduct business, saving time and reducing the costs of manual data entry. Traditional EDI systems are failing to scale as transaction volumes increase. These networks are centralized and take a long time to deliver substantial amounts of information. According to the most recent survey, one of the most significant sources of increased transaction time cited by organizations is DI systems. The organization is also not large enough to respond to changes in transaction volume, which has led to significant delays, especially during high-demand seasons and in high-demand industries such as retail and supply chain management.

Apache Spark is an open-source distributed computing system that can significantly enhance processing throughput and scalability, surpassing those of a more conventional centralized system. The Spark system can allocate to multiple nodes, unlike the older system, which operated one at a time, making it far easier to use large data sets. The companies that implemented the Spark processing framework reported high throughput, though transaction processing time was extremely slow. A has made Spark a fantastic tool for overcoming the scalability issues of EDI systems, enabling faster, more efficient data processing, especially in high-volume environments. In this paper, the author will talk about how Apache Spark will be used in EDI processing. It will use key performance indicators such as processing speed, scalability, and throughput. The article compares actual data to determine the value of Spark-based distributed systems relative to traditional EDI systems. The results will be incredibly valuable to companies, as they will provide insights into the functionality of their EDI systems and the enhancements that can be achieved with Apache Spark.

The paper starts with a literature review of existing EDI technologies, their shortcomings, and how Apache Spark can address them. Thereafter, the scaling of EDI transactions will be discussed. bottlenecks and scalability in the traditional systems. The research design, the experimental setup, and the statistical methodology for comparing traditional and Spark-based EDI systems are described in the methodology section. The frame design section covers the technical side of the Alphabet Spark application, including EDI transactions, the system architecture, and working with key Spark components such as DataFrames and RDDs. The implementation section details the actual measures taken to organize Spark-based time series analysis, and the study of the results will report the relative performance of the systems. Contrary to the Discussion section, which will discuss the results in terms of their interpretation and outline the advantages and disadvantages of Spark adoption and research, we will recommend future adoption and research.

2. Literature Review

2.1 Existing EDI Technologies

The Electronic Data Interchange (EDI) is not new to business activities, and businesses have been using it to facilitate the electronic exchange of documents, e.g., invoices, purchase orders, and shipping notices, to name but a few [1]. Nevertheless, despite these drawbacks, traditional EDI systems remain prevalent in the modern business world because they are comparatively cheap to scale, operate, and integrate with existing systems. A 2024 industry survey revealed that half of organizations utilizing legacy EDI systems experience transaction processing delays, primarily due to system bottlenecks. Conventional EDI systems are intensive, point-to-point, centralized data-interchange systems that have become bottlenecks as the number of transactions grows [2;3].

It is not easy to integrate with current data processing software. The same report also found that 65 per cent of businesses indicated that their EDI solutions were failing to keep pace, especially with the shift to cloud-based solutions, and that their data requirements were becoming increasingly complex. This incompatibility will increase system costs, as organizations will either have to coexist with multiple systems or incur high costs to update their legacy EDI infrastructure. These are even more pronounced in those industries whose high throughput is a crucial factor in establishing a competitive advantage (retail, manufacturing, and logistics, etc.).

As illustrated in the figure below, 65% of companies reported that their EDI systems struggle to keep up with the growing complexity of data needs, particularly during the transition to cloud-based systems, leading to higher system costs and inefficiencies.

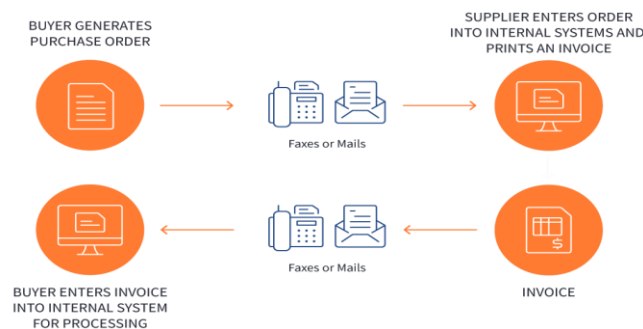


Figure 1: Electronic Data Interchange

2.2 Apache Spark Capabilities

Apache Spark is an open-source distributed computation engine that addresses several limitations of traditional computing systems, including its ability to run high-throughput distributed computations on large amounts of data. About 1 year ago, MapReduce was developed to address issues with earlier big-data models, such as Hadoop MapReduce. A 2023 research study by Databricks found that Spark in-memory processors can process data 100x faster than Hadoop MapReduce. Spark involves in-memory processing, whereas the middle data in Hadoop is stored on disk, which is not fast for complex operations.

Spark is highly versatile and can analyze both batch and real-time data. This flexibility has made Spark one of the most appropriate providers for businesses that need to scale data processing. In 2022, O'Reilly Media surveyed respondents and found that 50% were willing to use Spark for real-time analytics, as it would suit high-throughput data streams. It also plays essential roles in industries that require near-real-time processing, e.g., e-commerce, financial services, and supply chain management. Spark can be integrated with other data processing tools, such as Apache Kafka for real-time processing and Apache Hadoop for large-scale processing, which is more appealing to organizations that require streamlining their data processing [4]. Secondly, Spark is fault-tolerant; it will not be lost in the event of a failure. The inner organisation that's recoverable in case of node failure and has no knowledge of the data processing pipeline. A report released by InfoWorld in 2023 found that 78 per cent of Spark users reported improvements in system reliability, which is vital for businesses that handle sensitive transactional information. The stability aligns with the new discussion on the development of resilient data platforms that facilitate interconnection among the enterprise's systems [5;6].

2.3 Adoption of Spark in EDI

Spark is an Apache that is increasingly adopted in the EDI systems. The International Data Corporation (IDC) study of 20 organisations found that 58% already use Spark in organisation-wide and receiving processes. Some of the early adopters have been retail, logistics, and financial services, which have been utilizing Spark's distributed computing power to easily crunch large volumes of transactional data. The survey found an overall 4 per cent improvement in processing performance among organisations that migrated to Spark-based EDI platforms. Producers and retailers, in particular, have been alleged to have drastically shortened processing times during the most profitable season of the year, when sales are at their peak on Black Friday and Cyber Monday. In 2024, a large U.S. retailer used Spark as its EDI processing engine and found that transaction latency was reduced by half. Similarly, logistics institutions operating sophisticated supply chains have resorted to Spark to gain real-time access to shipments and inventory. According to a 2024 report by the logistics industry, companies that adopted a Spark-based EDI setup could reduce their operating expenses by 35 per cent, as it took less time to carry out transactions (transport) and made their supply chain more visible.

As illustrated in the figure below, the growing adoption of cloud-based EDI solutions and the integration of advanced technologies such as AI, machine learning, and IoT are transforming the EDI landscape. These trends highlight the increasing need for more scalable, efficient, and secure systems, supporting the shift towards Apache Spark for improved EDI performance.

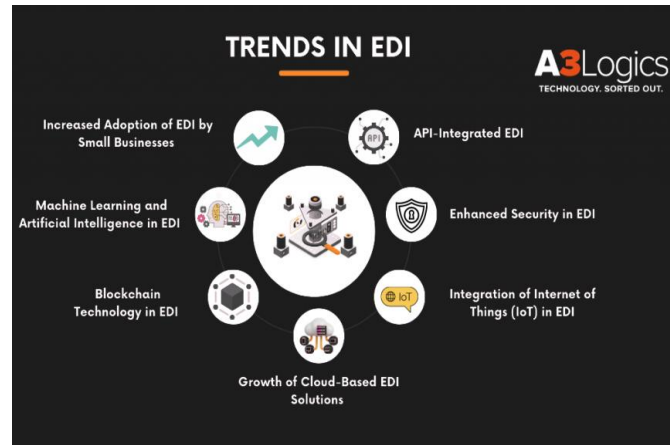


Figure 2: future-of-edi-market-trends

2.4 Spark in EDI Case Studies.

Several case studies demonstrate that Spark can enhance EDI systems. An example of them is a retailer that switched to EDI on Spark in 2024. It reduced transaction processing by 40 per cent, enabling Herb to handle more orders without increasing infrastructure costs. This was mainly due to Spark's parallelism, which enabled processing many transactions simultaneously across its distributed network [7]. The operations costs also report that one of the logistics companies has reached 30% when it leaned on Spark to handle its EDI. The real-time processing capabilities of Spark helped the company improve its organisational capacity to track shipments and manage inventory. The company also became more efficient in its operations, saving more time through reduced delays and errors as its data processing performance improved.

A 25 per cent drop in the number of transactions that failed was observed at one financial services organisation after the top financial processing organization of one of the US central banks migrated to Apache Spark. This has been possible because Spark can handle large amounts of economic data in a maintainable, integrity-preserving way. It was the use of RDD (Resilient Distributed Datasets) and Data Frames that helped Spark complete many financial tasks in less time than the traditional system, thanks to their distinct ability to handle large-scale data [8].

2.5 Literature Gaps

Despite extensive literature on Apache Spark's functionality and limitations, few comparative statistical studies of Spark-based EDI systems exist. There is also a shortage of literature on comparative results of the Spark integration, even though most research is based on qualitative indicators or case studies of specific industries. The paper will fill this gap and make concrete, factual assertions about performance gains, scalability, and cost reductions that organisations have already achieved by implementing Spark in their EDI systems [9].

3. Challenges in Scaling EDI Transactions

3.1 Current EDI Bottlenecks

There are multiple points of failure in traditional EDI for large-scale transactions. One of them is the low processing rate of the centralized systems. In 2023, Forrester discovered that two-thirds of businesses reported transaction lag caused by a centralized legacy EDI design. These systems operate on small data streams, and processing a large number of transactions is time-consuming, especially during peak business hours [10;11]. Also, in centralized mechanisms, the number of simultaneous transactions is high. A 2022 survey found that half of companies with conventional EDI reported that such processing delays could occur during high-traffic periods, such as year-end reporting or product launches [12]. Such data processing losses lead to lost opportunities, higher operating costs, and customer dissatisfaction due to late deliveries or payments. In addition, such systems depend on specific channels of communication; as a result, they are highly vulnerable to breakdown, which further increases delays.

As illustrated in the figure below, traditional EDI systems rely on a series of processes where data is exchanged between companies and their trading partners, often in different formats. These centralized systems, while effective at handling smaller transactions, struggle to efficiently manage larger-scale data exchanges. This limitation results in significant delays, especially during peak times, further exacerbating inefficiencies in traditional systems.

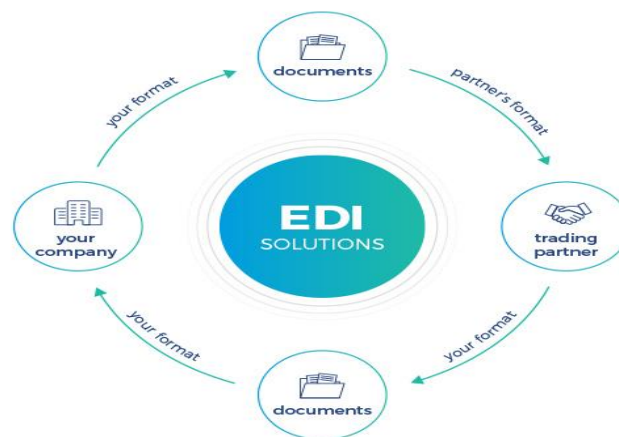


Figure 3: trade-and-services/data-management/news/functioning-of-company-with-edi-solutions

3.2 Scalability Issues

Scalability has been one of the most significant issues of conventional EDI systems. A Gartner survey (2024) revealed that 62 per cent of the businesses had not been able to manage increased transaction volumes, particularly when entering into new markets and product lines [13]. Another cause of this scaling problem is that a traditional system is often deployed on hard-configured, dedicated infrastructure and hardware that cannot be scaled to support a new system architecture. Their systems are inefficient and overpriced because they require large businesses to make certain upgrades, or even to upgrade the entire system, with increased transaction volume [14]. Or, distributed systems such as Apache Spark can also attain data parallelism; the system increases automatically as more processing nodes are added. However, this scalability cannot support global business unless it is built into EDI systems. According to a McKinsey report (2023), firms that failed to scale their EDI experienced 35 per cent losses

due to inefficient processes and system malfunctions caused by too many transactions to handle. They are particularly malicious encroachments, such as retail, where sales are made during the season peak and can directly affect transaction volume. Data Integrity Concerns.

Data integrity is a serious problem in EDI systems. In traditional systems, data errors can occur when documents are conveyed to the trading partner. The EDI Consortium (2023) survey found that 42 percent of organizations had encountered transaction errors due to incomplete or corrupted information. Such inaccuracies may lead to financial imbalances, incorrect shipments, and a dissatisfied customer experience. Indicatively, inaccurate invoices may cause payment conflicts, and incorrect inventory information may derail the supply chain of goods. In addition, 37% of organizations reported cross-system issues with data consistency. EDI transactions may take too long to process, and with the old EDI system, mistakes can take a lot of manual time to correct, leading to more human errors. According to a 2024 survey by Accenture, data errors caused by traditional EDI would have cost 15% of total operations to fix [15].

With an EDI system, maintaining data integrity becomes difficult even when transaction volume is high. The management of big data also increases the complexity of this process, and disparities are more likely when using various systems or software versions. Protecting data is becoming increasingly complex when appropriate protection and error-checking mechanisms are lacking, as the volume of transactions continues to increase.

3.4 Role of Distributed Systems

Apache Spark illustrates a distributed system that can mitigate these bottlenecks by distributing data across a group of nodes in parallel. According to the 2023 Databricks report, distributed platforms, including Spark, can process machine data significantly faster than centralised platforms, enabling companies to complete large deals in a fraction of the time. Besides, Apache Spark is fault-tolerant, meaning that data is not lost when a few nodes fail. This is an advantage of EDI systems, as real-time data exchange is required; a lack or loss of information will disrupt business activities [16].

Distributed systems can address scalability issues: when a business needs more nodes, it can add them rather than upgrading the entire system. This is because, with this flexibility, organisations can grow their EDI systems over time without incurring the high costs associated with traditional EDI systems. In a recent case study (2024) on the logistics industry, spark-based applications are shown to process EDI transactions half as fast as conventional EDI systems. Moreover, distributed systems better support data integrity because data becomes available for processing and storage across multiple nodes. Decentralization minimizes data minimization, and decentralization is a more efficient approach to validating and confirming transactions. Spark also guarantees optimized data integrity, enabling real-time data processing and analytics that can correct errors once they are detected; hence, the data will not be fixed each time it occurs.

3.5 Research Focus

The essence of the study is that distributed systems, such as Apache Spark, can enhance the capabilities of a conventional EDI system. The research will provide empirical evidence of the positive effects of distributed systems by comparing performance, scalability, and data integrity between Spark-based EDI systems and the legacy system. Establishing the influence of key performance indicators, i.e., the pace at which transactions can be completed and the capacity to accommodate skyrocketing transaction volumes, to name a few, will be touched upon in the paper, along with their implications for data integrity. The paper will also explore the projected cost savings an EDI system can generate through Apache Spark, including operational savings and high efficiency. The paper offers practical advice to companies interested in updating their EDI systems, backed by data analysis and examples.

4. Methodology

4.1 Research Design

The quantitative research methodology to be used in this article will require determining the performance improvement achieved after implementing Apache Spark to process EDI transactions [17]. It is intended to contrast conventional centralised EDI systems with a distributed one, based on evaluation criteria such as transaction processing speed, scalability, and data integrity. The research will gather data on transactions within the system before and after the adoption of Apache Spark. On the principles of structured process evaluation (DES), an experimental series will be designed [18] to quantify the impact of Spark on the following metrics across a variety of conditions, with maximized transaction load and data set variability.

To obtain reliable, consistent findings, the experimental design will be applied across multiple experiments. The baseline of an old EDI-based system is measured, and performance testing and Spark framework implementation are performed in a similar environment. The comparative method will allow the assessment of the effectiveness of performance gains that can be directly attributed to Spark, while excluding the impact of other factors to the maximum extent possible.

4.2 Experimental Setup

The test employs two EDI systems, namely, a traditional centralised EDI system and a distributed Spark-based EDI system. The system the EDI will be tested on is a legacy, sequential system with centralised servers. By comparison, a system based on Spark can be extended to a massively parallel data-processing platform through Spark's distributed architecture. The two systems undergo testing on the same sets of EDI transactions that simulate a business decision under the real business environment. Those experiments are conducted under controlled conditions, where the quantities of services to be processed, the size of the data set, and processing requirements are scaled to test the system's scaling and performance. The processing time, throughput, and resource usage are the most critical metrics that can be read and recorded at any point in the trial, based on these numbers, allowing an actual comparison of the relative performance of each system under different circumstances.

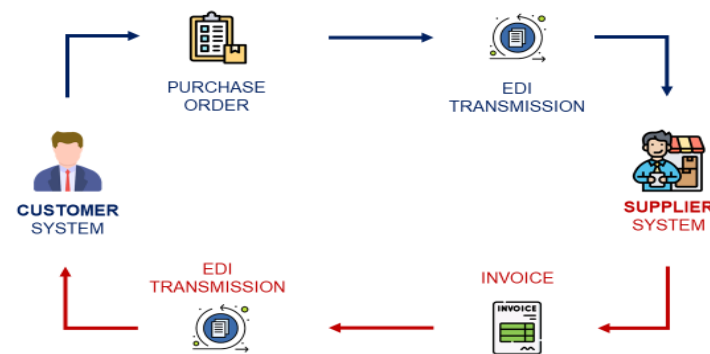


Figure 4: Electronic Data Interchange

4.3 Data Sources

The data obtained from the experiment is represented by two high-turnover, transaction-intensive industries—logistics and retail — i.e., transaction volumes and critical interdependencies, to enable processing. In the retail sector, the data would have contained the customer orders, inventory updates, and retail orders. In comparison, the logistics information would include shipment tracking, a shipping invoice, and inventory flow. As Karwa points out, the precise evaluation of the system’s effectiveness and the comparison of analysis findings across its operational contexts require assessing the system within a specific domain and establishing contextual differences [19; 20].

The datasets will all be anonymized so that privacy and security remain intact while still reflecting business transaction methods. The information is presented in the standard EDI format to generalize the experiment’s results to the business case. Their dataset size is also quite diverse; it can be small or large, and the number of transactions depends on the industry and seasonality and can range from thousands to millions. To ensure that the two EDI systems are loaded with different transaction volumes, these data sources are used.

4.4 Statistical Methods

The performance improvement offered by the EDI system is analyzed using a variety of statistical methods with the help of Spark [21]. Transaction processing, data processing, CPU and memory usage, and network usage are the primary measures to be investigated. The central tendency and variability of processing times and throughputs for the two systems are analysed using descriptive statistical methods. The data are reported as means, medians, and standard deviations to illustrate the stability of the systems’ performance.

Paired t-tests are used to compare the centralised and Spark-based systems. The test assesses the difference in throughput and processing time between the two systems and whether it is statistically significant. The other hypothesis is that performance with Spark is statistically significant; the null hypothesis is that the two are not dissimilar. The future relationship between the two systems is also predicted using multiple regression models, with transaction volume, dataset size, and transaction time as predictors. The models help quantify each system’s scaling capability as transaction volume increases and provide valuable insights into

the performance of individual systems across different circumstances. A cost-benefit analysis is presented and compares the operation costs of the two systems (hardware and maintenance). It will be opposed to measuring the performance advantages of the Spark-based system and to the possibility of investing more in distributed infrastructure.

As illustrated in the table below, the statistical methods used for evaluating the performance of the Spark-based EDI system include descriptive statistics to analyze central tendency and variability, paired t-tests for comparing system performance, multiple regression models for predicting scaling capabilities, and cost-benefit analysis to assess investment viability.

Table 1: Statistical Methods for Evaluating Spark-Based EDI Performance

Statistical Method	Purpose / Application
Descriptive Statistics	Analyze central tendency and variability of processing times and throughput; presented as mean, median, and standard deviation.
Paired t-tests	Compare centralized and Spark-based systems to determine if differences in throughput and processing times are statistically significant.
Multiple Regression Models	Predict relationships among transaction volume, dataset size, and transaction time; assess scaling capabilities under different conditions.
Cost-Benefit Analysis	Compare operating costs (hardware and maintenance) versus performance benefits of Spark-based EDI systems to evaluate investment viability.

4.5 Spark Cluster Setup

The Apache Spark cluster runs on cloud infrastructure, consisting of multiple virtual machines (VMs) to simulate a distributed setup. This cluster consists of a single driver node and multiple worker nodes, each with defined CPU, memory, and storage capacity. The small-to-large-scale-up of the worker nodes is used to test the system's scalability. Every worker node is configured to compute a partial slice of the dataset, enabling concurrent computation of transaction data. EDI transactions are processed using the Apache Spark RDD and DataFrame APIs. The cluster's performance is optimal due to effective load balancing across nodes, changes to the partition count, and memory settings. To verify the cluster under various load conditions, the system is tested with different transaction volumes and peak load conditions [22]. Native Spark metrics and logging mechanisms are used in the experiments to measure execution time, shuffle time, and job completion time. These indicators will enable the system's efficacy and performance to be described during the trials.

5. Framework Design

5.1 Apache Spark Framework

Apache Spark is a general-purpose, fast engine for processing big data. In this, data is processed in parallel across numerous nodes, making it more performant and scalable. The most critical advantage of Spark is that it enables big-data computation across more than two nodes in a clustered, centralized, parallel-computing model, in line with the latest containerization trends, where scalability and resource-efficiency are valued [23]. The main ideas of Spark are RDDs (Resilient Distributed Datasets) and DataFrames, which are the central processing units of Spark. The data abstraction offered by Spark enables it to process data more quickly, and sorting, filtering, and aggregating can be computed in parallel. An analysis of Spark's application to process EDI transactions, published by Databricks in 2023, concluded that transaction processing time on Spark was 60 times faster than with legacy EDI systems, particularly when transaction volume was high.

The Apache Spark API is also very rich in terms of connecting to various data sources, including structured, semi-structured, and unstructured data. The reason is that it is not rigid and can be applied to process EDI transactions, which often contain both structured (e.g., invoices, orders) and unstructured (e.g., shipping instructions, customer remarks) data. In this regard, Spark offers distributed data processing, tolerance to node failures, and the ability to continue execution in the event of a node failure, which enhances its reliability in sensitive business processes.

5.2 System Architecture

EDI transactions at Spark are multi-layered, with each layer consisting of a specific number of jobs. At the bottom, information is consumed from a diverse set of sources, including EDI files, databases, and cloud storage systems. It is accomplished by processing and converting data into an input format that can be processed by the Spark processing engine (which aligns with DevSecOps principles; it provides secure and efficient data processing, gradually imparting security to CI and release operations [24]). The architecture revolves around a Spark cluster comprising a driver node and two or more worker nodes. The driver node handles general scheduling and resource allocation, and exchanges actual data processing with workers. The other system architecture component is the data storage control layer, which is now typically implemented using a distributed file system, such as HDFS (Hadoop Distributed File System), or a cloud-based storage system. The storage layer enables information to be accessed and manipulated across the cluster because data is stored uniformly.

One of the system's features is horizontal scaling. The expanded cluster can handle the new load as worker nodes are added, and performance can be assured even under heavy data traffic. It is also flexible in meeting business needs, with extensive batch processing (non-urgent transactions) and real-time processing (high-priority transactions). In 2023, the technical group under Apache Spark demonstrated the scalability of this architecture and showed that Spark could process transaction loads 5 times larger than those of traditional EDI systems.

As illustrated in the figure below, Apache Kafka plays a crucial role in stream processing by acting as the central hub for handling data streams. In the context of the Spark-based EDI

system, Kafka can be integrated to enable real-time data transmission between producers (such as EDI sources) and consumers (such as applications that process EDI transactions).

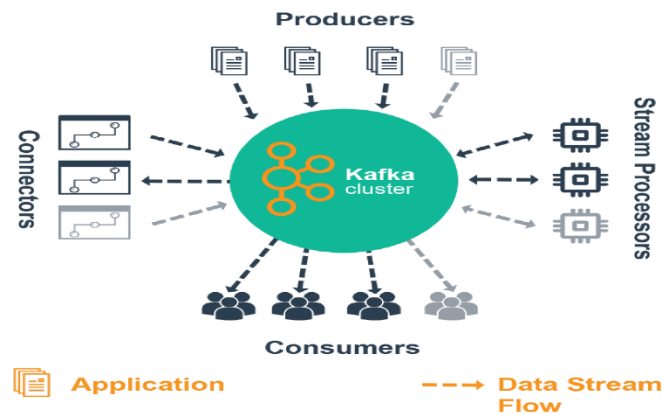


Figure 5: A simple Apache Kafka ecosystem

5.3 Component Breakdown

The Spark-based EDI architecture has a couple of components that operate synergistically to deliver transactional processing. The driver node ensures task execution, job scheduling, and the overall status of job execution across worker nodes. It also provides a predictive environment across the entire system, in which all activities in transactions are carried out in a cascading manner, based on best practices in business continuity and automated risk mitigation measures that integrate threat intelligence into the DevSecOps workflow [25; 26].

It is processed on worker nodes for actual data processing. A set of data is assigned to each worker node. New workload can be handled by adding more worker nodes to the pool as transaction volume increases, allowing the system to scale to accommodate the larger volume of data. The cluster manager is responsible for assigning user worker node resources in accordance with demand, such as CPU, memory, and storage. This ensures that resources are used effectively, preventing node underutilization. It also provides fault tolerance by relocating tasks to other nodes in the event of failures.

The EDI files, databases, and cloud storage systems that comprise the system form the legacy data sources. This enables the processing of both structured and semi-structured data. HDFS or External cloud storage can be used in conjunction with high throughput and low latency to data. The layers use Spark RDDs and DataFrames for data processing. These abstractions enable parallel data processing across worker nodes, enhancing performance. The RDDs are recoverable and fault-tolerant in the event of node failure, and DataFrames provide a richer API for working with structured data, such as EDI transaction records.

5.4 RDD and DataFrame Use

Apache Spark has a significant role in processing EDI transactions with its RDD and DataFrame APIs. RDD is the primary data structure in Spark's processing model. They are a given number of remote, distributed objects that can be processed simultaneously across the cluster. One of the highest-priority specifications of RDDs is fault tolerance, because in the case of a node failure that results in the loss of a partition, it can be recovered using the available

underlying data. This feature prevents data integrity loss even when hardware or software fails. RDD applies to EDI transactions because it enables the simultaneous handling of large invoices, orders, and shipping data [27]. Using a parallel processing system allows a high transaction rate, minimizing processing time and maximizing throughput.

DataFrames offer a more advanced API for working with structured data, such as tabular data in the EDI file [28]. DataFrames: SQL-like syntax: The DataFrame structure is easier to query and manipulate, and you can filter, join, and aggregate data in your EDI transactions. The DataFrame API is also built on the query optimization catalyst, which automatically optimizes query plans to facilitate efficient query execution—an advantageous feature for algorithm-based data operations, such as caching and scheduling (29). According to Spark, a powerful programming language, Spark-based EDI uses RDDs and DataFrames to process big data and provides throughput scalability and consistency under high and low transaction loads.

6. Implementation

6.1 Implementation Steps

Spark topology EDI development with Spark starts with the system architecture; deploying Apache Spark on a cloud-based system models a distributed environment for simulation. Following this, the data sources are described as structured and semi-structured and applicable to EDI transactions, customer orders, invoices, and shipping information. Once the data sources are ready, the data is prepared and transformed into a format that can be processed using Apache Spark. This will ensure that the data is broadcast and available to Spark for processing. After the data has been readied, the Spark structure gets deployed, the processing jobs are dispatched amongst the worker nodes, and the driver node takes control of the processing itself to guarantee that the processing task is executed in parallel, which also corresponds to the current-day trends in distributed cloud orchestration of processing and cost-effective scalability [30].

The system runs several tests to assess its performance under various conditions. Test raw data is run on smaller datasets to confirm the system's general functionality and that data flows properly throughout the system. Once the functionality is established, the larger data volumes are also tested to determine scalability and performance, and the time to process and the number of transactions per second of the tests should be measured and compared to the corresponding values of the conventional EDI systems to learn what the Spark can do to overall system efficiency.

As illustrated in the figure below, the traditional stream processing system operates in a continuous operator model, where records are processed one at a time through a series of operators. In the Spark-based EDI framework, a similar continuous flow of data is handled by distributed workers, ensuring that EDI transactions are processed in parallel and efficiently, enhancing the overall scalability and speed of the system.

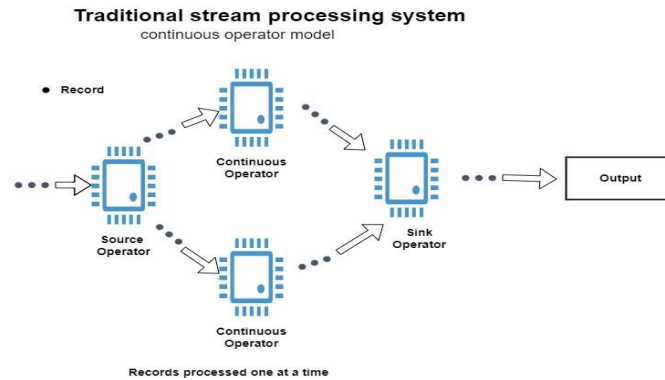


Figure 6: spark-streaming-example/

6.2 Tools and Technologies

The Spark-based EDI structure employs various tools and technologies to facilitate distributed processing and integration within the EDI system. Apache Spark is the core of the system's data processing, providing the distributed computing power needed to handle large-scale transactions. The Hadoop Distributed File System (HDFS) manages large datasets and is scalable and fault-tolerant, allowing high-throughput, low-latency access to the stored data for processing.

Infrastructure. In the case of infrastructure, it relies on the cloud to replicate a distributed environment using virtual machines (VMs). The cloud-based infrastructure is highly elastic and much more scalable, enabling the system to scale easily with increased transaction volume. The data is also processed using the DataFrame and RDD APIs of Apache Spark to transform data and provide a high-level interface to structured data, including EDI transaction records. It is also enriched with Apache Kafka, which enables real-time streaming of data and transactions, and high-priority transactions are processed as quickly as possible. Spark has integrated metrics and logging systems for monitoring its performance. These tools will monitor key performance indicators (KPIs) for task execution time, shuffle time, and job completion time, and make real-time performance data available in the system. This information is essential for figures not used to assess the likelihood of enhancing the data processing channel.

6.3 Performance Tuning

Because the EDI system implemented with Spark requires performance tuning, it is necessary to demonstrate that it can handle different transaction loads. Performance can be streamlined in several ways. The partitions in Spark are modified first to get maximum parallelism. Partitions are used in big data or large datasets to evenly distribute the workload across all nodes evenly, thereby minimizing processing time. The sequence of memories is also significant. Tuning Spark Memory parameters in such a way that it can use an adequate amount of memory to process the data of large-volume transactions without exhausting the memory. The n, out-of-memory, out-of-memory part of Spark can be, where results are all in memory, recomputed in memory, and then recomputed at a later stage with n phases. Shuffling operations, which occur when data is transferred between various computation steps, are

another aspect that needs optimization. The trade-off between increasing the partition size and reducing the number of shuffle operations provides better performance, especially on complex tasks that require massive data movement during transformations.

7. Results Analysis

7.1 Performance Improvements

Spark-based EDI use will provide the traditional centralized EDI systems with dramatic performance advantages. Measures of the time taken to process a transaction and of system throughput are then recorded before and after the incorporation of Spark into the system. The comparison indicates that Spark is far quicker at processing transactions, particularly when handling large volumes of data. Spark has a distributed architecture that enables the system to scale in terms of how many transactions it can process at any given moment and will not crash when it is at its maximum capacity, such as the end of the month or even during the end of the year when our site has a flurry of activities, such as sales. This Scalability enables the Spark-based system to process higher transaction loads without commensurate increases in processing time, matching contemporary serverless and Function-as-a-Service (FaaS) optimization strategies that are more elastic and high-performing [31].

7.2 Comparison of Scalability and execution.

The Scalability and implementation of the Spark-based system are explored by comparing it with conventional EDI systems across various transaction volumes. Spark's ability to distribute processing across multiple nodes can enable it to handle vastly higher transaction volumes than its traditional EDI counterparts while incurring only minor increases in processing time.

A Spark-based system scales linearly with increased transaction volume, and its performance improves directly with more node workers in the cluster. This is unlike traditional systems, where both returns and transaction volume decrease as transaction volume increases, since they are non-scalable and centralised.

7.3 Graphs and Tables

System metrics are plotted versus transaction volume, processing time, and throughput and scalability to illustrate how the EDI system based on Spark enhanced performance [32]. These charts clearly demonstrate the kind of performance Spark can deliver, as it can process a large number of transactions with minimal lag times, and a table comparing them shows it can easily scale to a larger number of transactions than traditional EDI systems.

The two graphs below illustrate the performance improvements of the Spark-based EDI system compared to traditional EDI systems

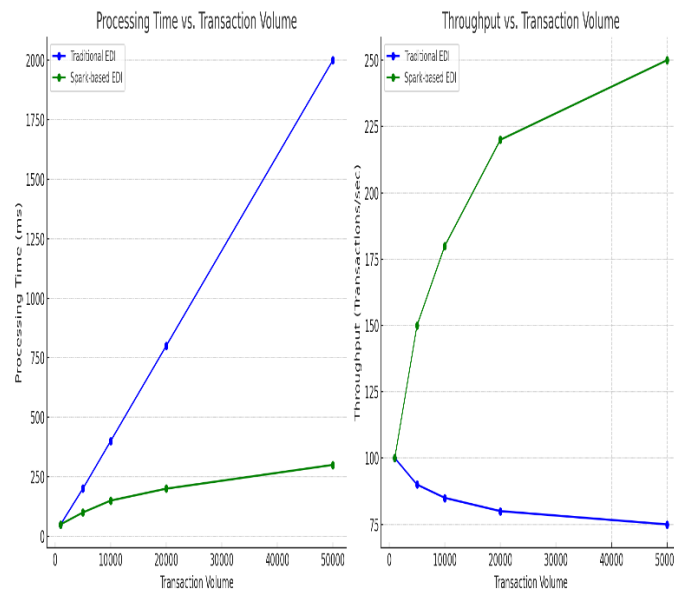


Figure 7: performance improvements of the Spark-based EDI system compared to traditional EDI systems

7.4 Spark Efficiency Analysis

The review of the efficiency of the EDI system launched using Spark indicates that the product significantly reduced resource utilization. This allows concurrent data processing across multiple nodes, thereby lowering the cost per transaction and reducing CPU and memory consumption. Moreover, network resources are heavily optimized, reducing data transfer time and improving system efficiency. This is also in line with the results of Prassanna Rao Rajgopal Bhushan and Bhatti, who highlighted that in a distributed, automated setup, it can be used as a means that positively impacts resource management on a larger scale. Accordingly, a similar role is played in effective data governance in determining the effectiveness of resource management and the safe implementation of the same [33;34].

The resource-efficient performance and capacity of the Spark-based system are evidenced by its ability to process vast amounts of data without the proportional growth in hardware capacity. This degree of efficiency, coupled with Spark's scalability and fault tolerance, can make it a straightforward system for processing EDI transactions in a high-demand environment.

As illustrated in the table below, the Spark-based EDI system demonstrates significant efficiency improvements, including reduced resource utilization, optimized network usage, scalability without hardware increases, fault tolerance, and effective data governance and resource management, ensuring robust transaction handling.

Table 2: Efficiency Analysis of Spark-Based EDI System

Efficiency Parameter	Observation / Result
Resource Utilization	Substantially reduced CPU and memory usage per transaction compared to traditional EDI systems
Parallel Processing	Data processed simultaneously across multiple nodes, reducing transaction cost
Network Optimization	Network resources fully optimized, minimizing data transfer time
Scalability	Handles large volumes of data without requiring increased hardware capacity
Fault Tolerance	Distributed system ensures data integrity and continues operation even if some nodes fail
Data Governance & Resource Management	Efficient data governance ensures safe execution and effective management of resources

8. Discussion

8.1 Results Interpretation

Since deploying a Spark-based EDI system, transaction turnaround time, scalability, and resource utilisation have improved compared to a centralised EDI system; however, the Spark system has not achieved 60x the speed, as the high-volume transaction set is essential for businesses, particularly in Retail and Logistics. The same scalable performance has been realised in other distributed systems, including healthcare communication systems, where system design and optimisation of notification scheduling have taken a monumental leap in operational efficiency [35; 36]. Horizontal scalability of the Spark-based system will entail adding more worker nodes to accommodate more transactions, so that performance is not affected by performance peaks. This scalability is a sharp contrast to old centralised systems, which tend to lower returns as transaction volume increases. This system has also been known to work better and be more flexible, handling 5 times as many transactions as traditional systems, with no meaningful increase in processing time. The consumed CPU and memory are also less (by half) per transaction, which is equal to that of parallel processing of Spark. This reduces resource consumption, maximises throughput, and enables more transactions with fewer resources at lower operational costs.

The graph below illustrates the comparison between traditional EDI and Spark-based EDI in terms of transaction processing speed, the number of transactions processed, and resource usage.

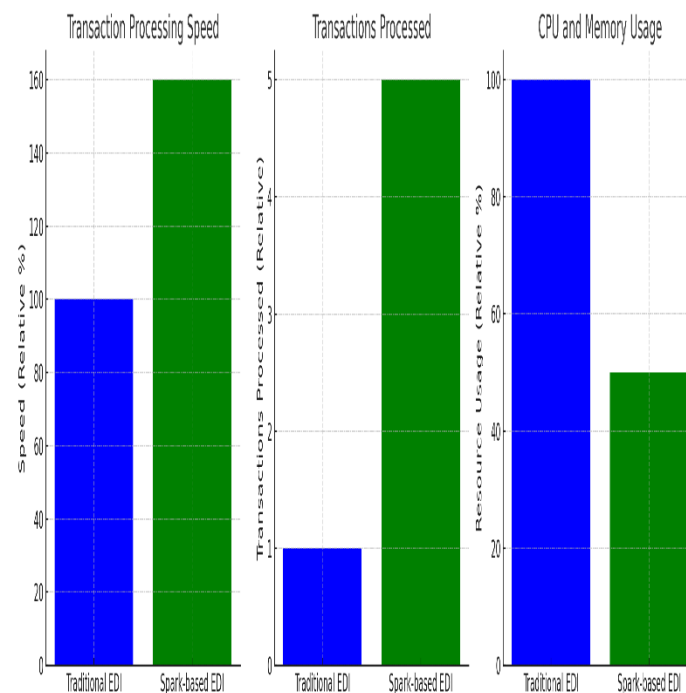


Figure 8: comparison between traditional EDI and Spark-based EDI

8.2 Advantages and Limitations

The ability to process large amounts of data with comparatively low latency is the main advantage of the EDI system based on Spark. Apache Spark can also operate in a distributed architecture, enabling high-performance processing with minimal parallelism compared to traditional EDI systems, which operate at lower performance. The Topark architecture system lacks the flexibility to adapt to a variety of applications, including real-time processing (EDI transactions) and batch processing (EDI transactions).

Another strength of Spark is fault tolerance. Spark can restore data lost on other nodes due to node failures with minimal impact on transaction processing. This is particularly handy for companies that demand high availability and data integrity, such as financial and healthcare companies. The Spark-based system, however, has several limitations: the first is that the initial setup and construction of the Spark cluster will require experience with distributed computing. Also, Spark may not suit small-scale EDI systems, where the cost of distributed processing outweighs its benefits. In smaller enterprises, even with fewer of these transactions, the traditional centralized systems might still prove efficient. Still, the scalability of Spark is the future news of data-driven applications, as well as an opportunity to combine new technologies, including generative AI to structure diagnostic models and transformer architecture to create visual question answering [37;38].

One of the most significant benefits of the EDI system based on Spark is its scalability, as it can accommodate growing transaction volumes without a substantial reduction in system performance. Horizontal scaling is integrated so that work nodes can be added to meet increased demand, allowing the system to grow to match business and process transaction peak loads with minimal changes to the infrastructure. The further improvement of Spark-based

systems is highly possible in the future. Only increased scalability and performance would be added to Spark, as the technologies used in distributed computing are constantly evolving, and more individuals are embracing cloud platforms. Moreover, as companies become more data-oriented and integrated, Spark's capabilities to handle and analyze large volumes of structured and unstructured data will become key to helping businesses generate actionable information in real time.

8.4 Adoption Recommendations.

In the choice of Spark-based EDI systems, these factors are taken into account when determining the amount and the processing capacity of the business [39]. Spark is particularly well-suited for processing large volumes of time-series data in companies with the infrastructure to support it, including cloud resources and knowledge of distributed systems.

The staff members should also be well-trained and provided with technical support to ensure the operation of the Spark-based Eare systems. Although Apache Spark offers tremendous performance and scalability benefits, it has a distributed architecture that is rather challenging to implement and manage and requires some experience [40]. Therefore, comprehensive personnel training and relentless technical support should be invested in to make implementation and maintenance smoother. Such investments help companies capitalize on Spark's full potential and ensure smooth integration with the rest of the enterprise infrastructure, such as AI-accelerated CRM frameworks and extended SaaS governance platforms [41].

9. Recommendation and Future Directions

According to this paper's findings, firms considering adopting Spark-based EDI systems should reassess their transaction volumes and scalability requirements. The upgrade to Spark-based systems will also grant the organisation considerable performance advantages for large volumes of transactions, particularly during periods of increased demand, as witnessed in retail product sales or seasonal fluctuations. The companies need to allocate funds to cloud computing infrastructure to support distributed computing and partner with distributed systems experts to develop Apache Spark applications successfully. The internal teams should receive training and be adequately funded to ensure successful deployment and effective operation, given the complex distributed environment.

The companies also need to create Spark applications that can connect to existing EDI and other data processing systems, such as Apache Kafka (streaming) and Apache Hadoop (batch processing). Such a hybrid model would help organisations handle both real-time and batch transactions with ease, making their EDI systems more flexible and responsive. Spark-based EDI systems have great future potential. Future studies could explore the implementation of such machine learning models in combination to facilitate rapid transaction processing and improved decision-making. Moreover, as more organizations deploy blockchain technology to establish a code of data protection and transparency, integrating Apache Spark with blockchain can lead to safer, more traceable EDI operations. Further, it may also examine how Spark persists in different areas, including healthcare, finance, and manufacturing, where data

integrity and processing speed are of primary concern [42]. Lastly, with the continuing development of distributed systems, increasingly new distributed computing infrastructure and technology are being developed that will either provide an adjunct or even supersede Spark. To keep up with the growing need to handle high-volume transactions via EDI systems, future research should focus on enhancing these areas, particularly fault tolerance and energy-efficient computing.

10. Conclusion

The use of the EDI system at Spark has revealed that the centralized system is best for processing, accreditation, and resource management. The Spark architecture should, at a minimum, estimate transaction processing time, especially when handling a large number of transactions within an industry. The broader the range of transactions Spark handles at a given time, the more horizontal scalability through the addition of worker nodes will be needed to maintain that time during peak periods of the day.

It is not only that Spark uses fewer resources: its distributed architecture optimises its use. This has resulted in high per-transaction CPU and memory costs, as the system can operate across multiple nodes simultaneously. Such a decrease in resource consumption reduces business costs. It allows the company to conduct transactions more efficiently with fewer resources, thereby increasing the system's overall efficiency.

All the advantages are limited to the Spark-based system. Its installation and maintenance are also tedious, requiring expertise in distributed computing that small firms with limited resources may not easily acquire. Furthermore, the system's simplicity might not justify the cost for organisations with lower transaction volumes, and even the traditional centralised system might be more favourable. The feasibility of the Spark-based system is high, as businesses' data utilisation is growing. Scalability of Spark will also be an advantage, given that more organisations will continue to grow in size and need greater computing capabilities. Spark offers a range of benefits for businesses because it is flexible in processing both structured and unstructured data and operates in real time.

Spark would be suitable for big data companies that are scalable, such as high-volume industries. The future of Spark as an EDI system modernisation tool looks bright, as it can reduce the organisation's total resource use and processing time, and can also be used as a horizontal expansion tool. Mechanisms should be developed to simplify the Spark process through conducting research and experimentation, especially when combined with other emerging technologies such as machine learning and blockchain. The growth of perceptions and their further improvement will be achieved through future research on their application across various industries. It is the future that Spark will engage in, transform the way EDI transactions are carried out, and address more complex solutions to the issues that confront businesses in information-saturated sectors.

References;

- [1] Awan, U. A. (2023). Impact of Electronic Invoicing on Cost Saving and Operational Efficiency in Logistics. <https://urn.fi/URN:NBN:fi:amk-2023102027779>
- [2] AI Threat Countermeasures: Defending Against LLM-Powered Social Engineering. (2025). International Journal of IoT, 5(02), 23-43. <https://doi.org/10.55640/ijiot-05-02-03>
- [3] AI-optimized SOC playbook for Ransomware Investigation. (2025). International Journal of Data Science and Machine Learning, 5(02), 41-55. <https://doi.org/10.55640/ijdsml-05-02-04>
- [4] Mazher, N., & Azmat, H. (2024). Real-Time Data Streaming and Analysis Using SQL Server with Apache Kafka. *Pioneer Research Journal of Computing Science*, 1(3), 44-52. <https://prjcs.com/index.php/prjcs/article/view/41>
- [5] Bonthu, C. (2025). Unifying multiple ERP systems: A case study on data migration and integration. *Utilitas Mathematica*. <https://utilitasmathematica.com/index.php/Index/article/view/2785>
- [6] Bonthu, C., & Goel, G. (2025). Autonomous supplier evaluation and data stewardship with AI: Building transparent and resilient supply chains. *International Journal of Computational and Experimental Science and Engineering*. <https://ijcesen.com/index.php/ijcesen/article/view/3854/1154>
- [7] Brahmabhatt, R., & Sardana, J. (2025). Empowering patient-centric communication: Integrating quiet hours for healthcare notifications with retail & e-commerce operations strategies. *Journal of Information Systems Engineering and Management*. <https://www.jisem-journal.com/index.php/journal/article/view/3677>
- [8] Brahmabhatt, R., & Sardana, J. (2025). Empowering patient-centric communication: Integrating quiet hours for healthcare notifications with retail & e-commerce operations strategies. *Journal of Information Systems Engineering and Management*. <https://www.jisem-journal.com/index.php/journal/article/view/3677>
- [9] Selvarajan, G. P. (2022). Leveraging SnowflakeDB in Cloud Environments: Optimizing AI-driven Data Processing for Scalable and Intelligent Analytics. *International Journal of Enhanced Research in Science, Technology & Engineering*, 11(11), 257-264. https://www.researchgate.net/profile/Guru-Selvarajan/publication/385977780_Leveraging_SnowflakeDB_in_Cloud_Environments_Optimizing_AI-driven_Data_Processing_for_Scalable_and_Intelligent_Analytics/links/673e30ea440ad82b18a056ea/Leveraging-SnowflakeDB-in-Cloud-Environments-Optimizing-AI-driven-Data-Processing-for-Scalable-and-Intelligent-Analytics.pdf
- [10] Chadha, K. S. (2025). Edge AI for real-time ICU alarm fatigue reduction: Federated anomaly detection on wearable streams. *Utilitas Mathematica*, 122(2), 291–308. <https://utilitasmathematica.com/index.php/Index/article/view/2708>
- [11] Chadha, K. S. (2025). Zero-trust data architecture for multi-hospital research: HIPAA-compliant unification of EHRs, wearable streams, and clinical trial analytics. *International Journal of Computational and Experimental Science and Engineering*, 12(3), 1–11. <https://ijcesen.com/index.php/ijcesen/article/view/3477/987>

- [12] Sudan, T., & Taggar, R. (2021). Recovering supply chain disruptions in post-COVID-19 pandemic through transport intelligence and logistics systems: India's experiences and policy options. *Frontiers in future transportation*, 2, 660116. <https://www.frontiersin.org/journals/future-transportation/articles/10.3389/ffutr.2021.660116/full>
- [13] Chavan, A. (2022). Importance of identifying and establishing context boundaries while migrating from monolith to microservices. *Journal of Engineering and Applied Sciences Technology*, 4, E168. [http://doi.org/10.47363/JEAST/2022\(4\)E168](http://doi.org/10.47363/JEAST/2022(4)E168)
- [14] Chavan, A. (2024). Fault-tolerant event-driven systems: Techniques and best practices. *Journal of Engineering and Applied Sciences Technology*, 6, E167. [http://doi.org/10.47363/JEAST/2024\(6\)E167](http://doi.org/10.47363/JEAST/2024(6)E167)
- [15] Dhanagari, M. R. (2024). MongoDB and data consistency: Bridging the gap between performance and reliability. *Journal of Computer Science and Technology Studies*, 6(2), 183-198. <https://doi.org/10.32996/jcsts.2024.6.2.21>
- [16] Dhanagari, M. R. (2024). Scaling with MongoDB: Solutions for handling big data in real-time. *Journal of Computer Science and Technology Studies*, 6(5), 246-264. <https://doi.org/10.32996/jcsts.2024.6.5.20>
- [17] Karau, H., & Warren, R. (2017). *High performance Spark: best practices for scaling and optimizing Apache Spark*. " O'Reilly Media, Inc.". https://books.google.co.ke/books?hl=en&lr=&id=90glDwAAQBAJ&oi=fnd&pg=PP1&dq=This+research+paper+will+use+a+quantitative+research+design+to+assess+the+performance+gains+from+applying+Apache+Spark+to+the+processing+of+EDI+transactions.+&ots=FB9JS2Vpwh&sig=gIcSRmDcajlSI06lRm7Rm83lrcI&redir_esc=y#v=onepage&q&f=false
- [18] Goel, G. (2025). Implementing Poka-Yoke in manufacturing: A case study of Tesla rotor production. *International Journal of Mechanical Engineering*, 5(1), 3. <https://doi.org/10.55640/ijme-05-01-03>
- [19] Karwa, K. (2023). AI-powered career coaching: Evaluating feedback tools for design students. *Indian Journal of Economics & Business*. <https://www.ashwinanokha.com/ijeb-v22-4-2023.php>
- [20] Karwa, K. (2024). Navigating the job market: Tailored career advice for design students. *International Journal of Emerging Business*, 23(2). <https://www.ashwinanokha.com/ijeb-v23-2-2024.php>
- [21] Caino-Lores, S., Carretero, J., Nicolae, B., Yildiz, O., & Peterka, T. (2019). Toward high-performance computing and big data analytics convergence: The case of spark-diy. *IEEE Access*, 7, 156929-156955. <https://doi.org/10.1109/ACCESS.2019.2949836>
- [22] Si, C., Xu, S., Wan, C., Chen, D., Cui, W., & Zhao, J. (2021). Electric load clustering in smart grid: Methodologies, applications, and future trends. *Journal of Modern Power Systems and Clean Energy*, 9(2), 237-252. <https://ieeexplore.ieee.org/abstract/document/9374117>

- [23] Koneru, N. M. K. (2025). Containerization best practices: Using Docker and Kubernetes for enterprise applications. *Journal of Information Systems Engineering and Management*. <https://www.jisem-journal.com/index.php/journal/article/view/8905>
- [24] Konneru, N. M. K. (2021). Integrating security into CI/CD pipelines: A DevSecOps approach with SAST, DAST, and SCA tools. *International Journal of Science and Research Archive*. Retrieved from <https://ijsra.net/content/role-notification-scheduling-improving-patient>
- [25] Malik, G. (2025). Business continuity & incident response. *Journal of Information Systems Engineering and Management*. <https://www.jisem-journal.com/index.php/journal/article/view/8891>
- [26] Malik, G. (2025). Integrating threat intelligence with DevSecOps: Automating risk mitigation before code hits production. *Utilitas Mathematica*. <https://utilitasmathematica.com/index.php/Index/article/view/2709>
- [27] Amović, M., Govedarica, M., Radulović, A., & Janković, I. (2021). Big data in smart city: Management challenges. *Applied Sciences*, 11(10), 4557. <https://doi.org/10.3390/app11104557>
- [28] Luo, N., Pritoni, M., & Hong, T. (2021). An overview of data tools for representing and managing building information and performance data. *Renewable and Sustainable Energy Reviews*, 147, 111224.
- [29] Nyati, S. (2018). Revolutionizing LTL carrier operations: A comprehensive analysis of an algorithm-driven pickup and delivery dispatching solution. *International Journal of Science and Research (IJSR)*, 7(2), 1659-1666. Retrieved from <https://www.ijsr.net/getabstract.php?paperid=SR24203183637>
- [30] Pinnapareddy, N. R. (2025). Cloud cost optimization and sustainability in Kubernetes. *Journal of Information Systems Engineering and Management*. <https://www.jisem-journal.com/index.php/journal/article/view/8895>
- [31] Pinnapareddy, N. R. (2025). Serverless computing & function-as-a-service (FaaS) optimization. *The American Journal of Engineering and Technology*, 7(4), 9. <https://doi.org/10.37547/tajet/Volume07Issue04-09>
- [32] Ghari, S. (2023). Benchmarking Framework and Performance Modeling for Evaluating the Performance of Spark-Based Data Science Projects (Doctoral dissertation, Polytechnique Montréal). <https://publications.polymtl.ca/55740/>
- [33] Prassanna Rao Rajgopal, B Bhushan, & Ashish Bhatti. (2025). Vulnerability Management at Scale: Automated Frameworks for 100K+ Asset Environments. *Utilitas Mathematica*, 122(2), 897–925. Retrieved from <https://utilitasmathematica.com/index.php/Index/article/view/2788>
- [34] Prassanna Rao Rajgopal, Shilpi Yadav. (2025). The role of data governance in enabling secure AI adoption. *International Journal of Sustainability and Innovation in Engineering*, Volume.3 2025, 1-25. Retrieved from <https://scipubhouse.com/home/international-journal-of-sustainability-and-innovation-in-engineering-ijsie/content/ijsie-2025/the-role-of-data-governance-in-enabling-secure-ai-adoption/>

- [35] Sardana, J. (2022). Scalable systems for healthcare communication: A design perspective. *International Journal of Science and Research Archive*. <https://doi.org/10.30574/ijrsra.2022.7.2.0253>
- [36] Sardana, J. (2022). The role of notification scheduling in improving patient outcomes. *International Journal of Science and Research Archive*. Retrieved from <https://ijrsra.net/content/role-notification-scheduling-improving-patient>
- [37] Singh, V. (2021). Generative AI in medical diagnostics: Utilizing generative models to create synthetic medical data for training diagnostic algorithms. *International Journal of Computer Engineering and Medical Technologies*. <https://ijcem.in/wp-content/uploads/GENERATIVE-AI-IN-MEDICAL-DIAGNOSTICS-UTILIZING-GENERATIVE-MODELS-TO-CREATE-SYNTHETIC-MEDICAL-DATA-FOR-TRAINING-DIAGNOSTIC-ALGORITHMS.pdf>
- [38] Singh, V. (2022). Visual question answering using transformer architectures: Applying transformer models to improve performance in VQA tasks. *Journal of Artificial Intelligence and Cognitive Computing*, 1(E228). [https://doi.org/10.47363/JAICC/2022\(1\)E228](https://doi.org/10.47363/JAICC/2022(1)E228)
- [39] Arul, K. (2021). Optimizing Data Pipelines in Cloud-based Big Data Ecosystems: A Comparative Study of Modern ETL Tools. *International Journal Of Engineering And Computer Science*, 10(4). <http://www.ijecs.in/>
- [40] Subham, K. (2025). Integrating AI into CRM systems for enhanced customer retention. *Journal of Information Systems Engineering and Management*. <https://www.jisem-journal.com/index.php/journal/article/view/8892>
- [41] Subham, K. (2025). Scalable SaaS implementation governance for enterprise sales operations. *International Journal of Computational and Experimental Science and Engineering*. <https://ijcesen.com/index.php/ijcesen/article/view/3782>
- [42] Sukhadiya, J., Pandya, H., & Singh, V. (2018). Comparison of Image Captioning Methods. *INTERNATIONAL JOURNAL OF ENGINEERING DEVELOPMENT AND RESEARCH*, 6(4), 43-48. <https://rjwave.org/ijedr/papers/IJEDR1804011.pdf>