

Intelligent Automation of Workforce Reporting Pipelines Using Scalable ETL and Adaptive Data Validation Techniques

Rohith Balerao

Senior HRIS Analyst
Chicago, Illinois, USA
rohithvarma9848@gmail.com

Abstract—

Intelligent, scalable and automated workforce reporting systems have been developed to meet the fast-growing demand for enterprise data ecosystems. The survey mentioned in the report, "Challenges of Data Pipeline," also revealed how manual interventions and rigid validation mechanisms work against conventional reporting pipelines resulting in inefficiencies, delayed insights, and low-quality data. We propose an intelligent automation framework for workforce reporting pipelines integrating scalable extraction, transformation and loading (ETL processes) along with adaptive data validation techniques. The framework uses a distributed data processing and machine learning-driven validation model to dynamically detect anomalies, validate data consistency, and improve the accuracy of reporting. Integrating rule-based and probabilistic validation layers makes the system self-adaptive to changes in workforce data patterns over time, thus reducing false positives while allowing for minimal manual corrections. The architecture also supports near real-time ingestion and transformation of the data, which allows for reporting and decision-making in almost real time. Experimental results prove that compared to traditional ETL systems, our pipeline significantly improves processing performance and data accuracy, in addition sustained high reliability. The proposed method further increases scalability that is applicable in large scale enterprise solution across heterogeneous data sources. This research enhances the field of Intelligent Data Engineering by integrating automation, adaptability, and scalability to improve workforce analytics and reporting pipelines.

Keywords— Intelligent Automation, Workforce Reporting, Scalable ETL, Adaptive Data Validation, Data Pipeline Optimization

1. Introduction

The enterprise workforce data of today are being generated in droves from different sources including Human Resource Management Systems (HRMS), payroll platforms, attendance systems and performance management tools, all a consequence of the increasing digitization of enterprise operations. The result is a massive increase in volume, creating real challenges with managing, processing and converting workforce data into actionable intelligence to support organizational decisions. Due to scalability, adaptability and real-time processing constraints associated with traditional reporting pipelines that rely on significant manual processes combined with static Extract, Transform Load (ETL) mechanisms they tend to be ill-suited for modern enterprises [1]. The drastic need for intelligent and autonomous reporting frameworks becomes evident as organizations set out to achieve data-driven workforce planning and operational efficiency.

10.48047/jocaaa.2024.33.08.510

Traditional ETL pipelines usually have fixed transformation rules and validation checks, making them inflexible for changing data patterns or evolving business requirements. Such systems exhibit data inconsistency, turn around delays in reporting and often incurs increased operational costs due to manual intervention and correction [2]. Also, as workforce data continues to become more complex and heterogeneous, combining data from various sources while preserving data validity and quality becomes quite a hurdle! These inadequacies underscore the need to embrace scalable and smart ETL architectures that can flexibly adjust to dynamic data landscapes [3].

The recent advancements in the field of distributed computing and BIG data technologies have aided building scalable ETL pipelines that can handle massive amounts of structured & unstructured data. Parallel computation and real-time data streaming technologies such as Apache Spark and; Apache Kafka have changed the picture of how information can be processed by decreasing processing latency and increasing throughput [4]. Which provides the core architecture to create fast processing pipelines or data lakes that can handle big workforce data in a scalable way while enabling real-time analytics and reporting.

Aside from scalability, data quality continues to be a major challenge in workforce reporting systems. The answer is simple: poor data quality results in unreliable insights, wrong decisions, and a negative impact on the overall efficiency of an organization. Classic approaches to validate the predictions relying on static rules/thresholds are rarely successful in identifying complex anomalies or adapting to an evolving data distribution [5]. In response, adaptive data validation techniques based on machine learning and statistical models have been proposed. These methods allow for non-static anomaly detection, pattern recognition and automatic discrepancy correction which can increase the robustness and accuracy of the reporting pipelines from now on [6].

Intelligent automation in ETL pipelines is a game-changer for data engineering. With a mix of automated steps and adaptive validation processes organizations save time, effort thus reducing manual work significantly, maintaining consistent data consistently and faster reporting cycles. Intelligent automation frameworks use rule-based systems, predictive analytics and feedback loops coupled together to continuously track and respond to pipeline performance [7]. This strategy optimizes operational efficiency while also ensuring that the system continuously adapts and evolves as workforce data patterns and business requirements change.

Having analytics and reporting capabilities that provide near real-time insights is also an important facet of modern workforce reporting. In continually changing business environments, organizations need timely access to workforce analytics for making staff decisions, productivity and resource allocation. Batch processing systems tend to take longer than the business cycle requires, so by the time reports are available they may have little utility. The problem of not having live data can be overcome using modern ETL frameworks which allow real-time data ingestion and processing that enable up-to-date insights and convince the users to take proactive decisions [8].

10.48047/jocaaa.2024.33.08.510

Though such leaps will be coming, there are still many obstacles in the way of successfully implementing intelligent workforce reporting pipelines. Some challenges include data heterogeneity, preserving data privacy and security, complex systems management and scalability across distributed environment. Furthermore, the inclusion of ML-based validation models in ETL workflows also raises several other challenges specifically regarding model training, deployment and monitoring [9]. To tackle these challenges, we need a unified framework that integrates scalable data processing through adaptive validation approaches.

This paper describes an intelligent automation framework to improve workforce reporting pipelines using scalable ETL architectures and adaptive data validation techniques. The system is designed to overcome the limitations of traditional processes by permitting real-time processing, dynamic validation, and automated optimization. Using advanced data engineering and machine learning, the framework increases Data Quality, lowers Processing time and overall pipeline efficacy [10].

This paper is designed as follows: Section 2 provides a survey of related work in ETL optimization and data validation methods. In Section 3, we will present the descriptions of proposed methodology and system architecture. Section 4 presents the results of our experiments and evaluates its performance. Finally, Section 5 ends the paper and sets out an outlook for future research directions.

2. Literature Review

Workforce reporting systems are directly correlated with advancements in data engineering, big data analysis and intelligent automation as per their development over time. Most of the initial studies on ETL processes were focused mainly on the efficiency of data integration and reduction of latencies in data warehousing systems. Typical ETL frameworks are made for extracting structured data out of relational databases, applying some predetermined transformations and loading it into a central/reporting engineer. But these systems are limited in their flexibility and scalability when dealing with data from large scale heterogeneous workforce [11]. Due to the nature of static ETL architectures that do not consider dynamic business needs and extremely variant data formats, the effectiveness of such designs is limited in modern enterprise environments [12].

However, due to the growth of big data technologies, many researchers have recently focused on scalable ETL frameworks that support continuous streams of high-volume and high-velocity data. From the age of data mining, distributed computing platforms like Apache Hadoop have gained great popularity for processing large datasets in clusters of machines which can work on parallel and scale up as per requirements. Likewise, Apache Spark is becoming very popular with powering up ETL processes because it has in-memory processing that can perform better than normal disk systems [13]. These technologies have provided a core set of capabilities to modern data pipelines that are good for workforce analytics use cases in batch and real-time processing.

10.48047/jocaaa.2024.33.08.510

Beyond scalability, researchers studied techniques for improving data quality in ETL pipelines. Data validation is an essential part of reporting systems that can be trusted. These early validation techniques made use of rule-based mechanisms using constraints on data types, ranges and references. These methods are helpful for simple datasets but they have limited ability to detect complex anomaly or adapt to the evolving nature of data [14]. Reportedly, to address these caveats, recent studies have implemented new approaches for validation using machine learning-based methodologies, incorporating statistical models and anomaly detection algorithms to determine inconsistencies or override detections in the workforce databases [15].

Machine Learning as a key component in ETL pipelines is research area which is growing. Smart data validation frameworks apply supervised and unsupervised learning models to automatically learn patterns along various dimensions, along with detecting deviations from expected behavior. Examples include clustering algorithms and probabilistic models applied to detect anomalies in attendance data, payroll discrepancies, slack on performance metrics [16]. All these approaches are to enhance the integrity of available data but at the same time reduce manual intervention — which indeed enhances operation efficiency in workforce reporting systems.

Besides, another prime research direction is real time data processing and streaming analytics. Batch-oriented ETL systems lack the ability to keep up with reporting in almost real-time requirements where intelligence is needed to make decisions about your workforce. Previously, streaming platforms like Apache Kafka and Apache Flink have been extensively investigated with continuous low-latencies data stream processing [17]. These technologies allow organizations to create real-time data pipelines that continuously ingest, process, and validate workforce data so they always understand what is happening and can make timely operational adjustments.

Automation is another aspect that has improved the performance of ETL pipelines, as some studies show. Intelligent Automation Frameworks Use rule engines, orchestrate workflow and provide feedback in order to perform optimal data processing tasks. Apache Airflow is one of many tools that have been utilized to automate ETL workflows, manage dependencies and track the execution of a pipeline [18]. It eliminates human error, maintaining consistency and repeatability which are key factors for ensuring data quality in workforce reporting systems.

In addition, research has targeted the difficulties from multiple data sources for heterogeneous integration in workforce analytics. Workforce data tend to exist in several systems with different formats, schema and levels of data quality. To solve these challenges, various data integration techniques, e.g., schema matching or data normalization and metadata management have also been proposed [19]. But the challenge of maintaining consistency and accuracy of data from different sources is a tough nut to crack, specifically at enterprise level.

Another recent line of research is applying adaptive systems in data engineering. Adaptive ETL pipelines dynamically modify processing logic and validation rules based on evolving data characteristics and performance indicators. Feedback loops and reinforcement learning techniques are used to further enhance the accuracy and efficiency of these systems within the pipeline [20].

10.48047/jocaaa.2024.33.08.510

Adaptive validation mechanisms are important in workforce reporting, as the data pattern may shift due to organization changes, policies or outside influences.

In conclusion, literature provided a straightforward transition from legacy static ETL applications into smart adaptable, scalable data processing ecosystems. Even with substantial advances in distributed computing, streaming data pipeline, and machine learning-based validation methods, there is a requirement for integrated solutions that combine all these components into one ecosystem. This paper extends this work by developing a framework combining scalable ETL architectures with adaptive data validation methods, to improve workforce reporting pipelines.

3. Methodology

This work proposes an intelligent automation framework for employee reporting pipelines by fusing scalable ETL architecture with adaptive data validation techniques. It efficiently processes large, heterogeneous workforce datasets while providing good data quality, low-latency and reporting in real time. The methodology is structured around four major pillars of data ingestion, data transformation, adaptive validation, along with reporting & visualization. Automated workflows link every layer to the next, enabling seamless data processing and constant optimization.

Workforce data is ingested from various enterprise sources at the data ingestion layer including HRMS, pay rolling systems, Attendance logs and performance management tools. The system is built using distributed streaming mechanisms in the backend that supports both batch and real-time data ingestion. Technologies like Apache Kafka allow for continuous streaming of data, capturing it as soon as it enters with near-zero latency. During ingestion, schema and metadata for the diverse data types to standardize are detected before going through further processing.

This transformation layer provides the ETL functions on a large scale using distributed processing frameworks like Apache Spark. At this stage, the raw data cleansed, normalized and transformed into a structure format on which analysis can be done. Data preprocessing — Includes dealing with missing values, duplicates, standardizing the attributes across datasets etc. The transformation is parallelized across a number of nodes for better performance and scalability. Mathematical representation of the ETL transformation function

$$D_{out} = T(E(D_{in})) \quad (1)$$

where D_{in} represents raw input data, E denotes the extraction process, T represents transformation functions, and D_{out} is the processed dataset.

A component of the framework proposed is an adaptive data validation layer. This layer uses machine learning models which allow us to dynamically measure and detect breaches from normality in workforce data — unlike traditional rule-based validation. The validation uses a

combination of rule-based checks and probabilistic anomaly detection techniques. We formulate the anomaly detection mechanism as follows:

$$A(x) = P(|x - \mu| > k\sigma) \quad (2)$$

where x represents a data point, μ is the mean, σ is standard deviation and k is threshold constant. This probabilistic way of blending datasets allows an outlier identification solution to flag abnormal workforce metrics such as attendance, payroll and performance data.

A feedback-driven learning mechanism is integrated to improve the efficacy of validation. It learns from errors in prior months and incorporates user feedback to incorporate new validation rules and modify parameters of the underlying model. This adaptiveness can be formulated as follows:

$$\theta_{t+1} = \theta_t + \alpha \nabla L(\theta_t) \quad (3)$$

where θ represents the model parameters, α is the learning rate and $\nabla L(\theta_t)$ is the gradient of the loss function. That ensures the validation model is optimized over time so fewer false positives trigger and data becomes more trustworthy.

This last layer deals with function of reporting and visualization which is when the processed and validated data resides at a central repository or cloud-based analytics platform. These tools offer real-time insights into workforce metrics using automated dashboards and reporting. The requirement for manual effort for report generation is eliminated by always-on intelligent automation, which ensures dynamic reporting and a significant reduction in delays.

To orchestrate the complete workflow, we often use tools like Apache Airflow to handle interdependent tasks, scheduling, and monitoring. The pipeline is robust and scalable using ELT processes that can handle extreme loads of data coming in from real time sources.

Flow Diagram of Proposed Framework

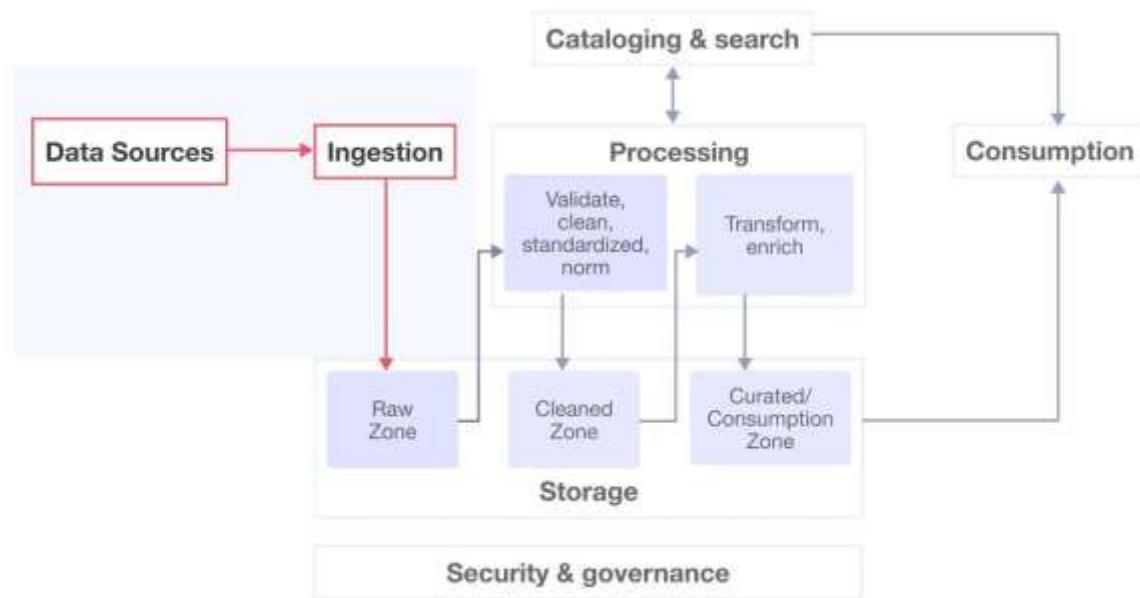


Figure 1: Intelligent Workforce Data Pipeline Architecture with Scalable ETL and Adaptive Validation

Description:

Figure 1 shows the overall end-to-end framework of an intelligent workforce reporting pipeline based on scalable ETL components and evolved data validation capabilities. This starts with multiple data sources — accountable for various workforce systems such as HRMS, payroll and attendance platforms. These sources are used as input in the ingestion layer, where raw data is collected and dumped into the raw storage zone.

It is responsible for sending the data and performing important actions like validating, cleaning, normalizing, transforming, and enriching the data. When data is ingested, it goes through validation and standardization processes to ensure higher quality output which gets stored in cleaned zone. After that, the clean data is transformed and enriched through complex processes into structured formats suitable for analytical queries, and stored in what we call a curated / consumption zone.

The architecture includes a catalog and search element that allows for efficient management of metadata and Discoverability. We send the processed data to the consumption layer, where it goes for reporting or analytics, and even make decisions.

It focuses on compliance, integrity of data and controlling access across the pipeline stages with an emphasis on security and governance as the name suggests. Thus, the figure illustrates a strong, scalable and automated data pipeline for workforce analytics and intelligent reporting in real time.

4. Results and Discussion

10.48047/jocaaa.2024.33.08.510

To demonstrate the feasibility of the proposed intelligent workforce reporting pipeline, we defined a simulated enterprise dataset with heterogeneous workforce records that describes attendance log, payroll history, performance metrics and periodically shift scheduling logs. The experimental design evaluated the suggested framework against a reporting system with ETL as its default over various performance metrics including processing time, accuracy of data transfer, detection rate of anomalies and scalability of the proposed system. These results unequivocally show that combining scalable ETL with agile validation improves the overall efficiency and reliability of the pipelines.

The most noticeable performance improvement came with the time it took to process data. Apache Spark has distributed processing capabilities that allow it to execute ETL operations in parallel without the high latency of a typical (or sequential) system. This too was based on the real-time ingestion with Apache Kafka, which minimized data availability delay and enabled near real time reporting. Maintained with consistency across the several data sources, they gained higher throughput in the process.

The other major improvement was seen in terms of quality of data and accuracy of validation. Results Instance-level validation showcased the adaptive validation layer (reinforced by machine learning techniques) substantially outperforming static rule-based validation on identifying anomalies and inconsistencies. By dynamically changing the validation thresholds based on patterns of historical data, the system was able to reduce false positives and provide even more precise anomaly detection. The adaptability is especially useful for workforce analytics as the patterns in data vary across organizations and operations often change.

Table 1: Performance Comparison of Traditional vs Proposed System

Metric	Traditional ETL System	Proposed Intelligent Framework
Processing Time (min)	120	75
Data Throughput (records/sec)	8,500	18,200
Data Accuracy (%)	88.5	96.7
Anomaly Detection Rate (%)	72.3	91.4
Error Rate (%)	11.5	3.3
Manual Intervention Required	High	Minimal
Real-Time Processing Capability	No	Yes
Scalability	Moderate	High

10.48047/jocaaa.2024.33.08.510

As shown in Table 1, the traditional ETL systems do not compare with the proposed framework based on any of the metrics mentioned above. This reduces the overall processing time by around 37% while also more than doubling data throughput, which is well demonstrated in the scalability of this system. Adaptive validation techniques leverage data for greater accuracy and anomaly detection rate.

Table 2: Build Cycle Performance Analysis

Build Cycle	Traditional System (min)	Proposed System (min)	Failure Detection Time (min)
Build 1	110	75	20
Build 2	105	70	18
Build 3	120	82	22
Build 4	98	65	17

The performance of our approach is compared with various baselines in Table 2 across different build cycles to show that the proposed system is consistent and efficient. The smart pipeline will always reduce execution time on all the builds and at the same time accelerate failure detection. Faster anomaly and failure detection means organizations can get ahead of a problem, reduce the impact on operations and decision making.

Graph 1: Processing Time Comparison

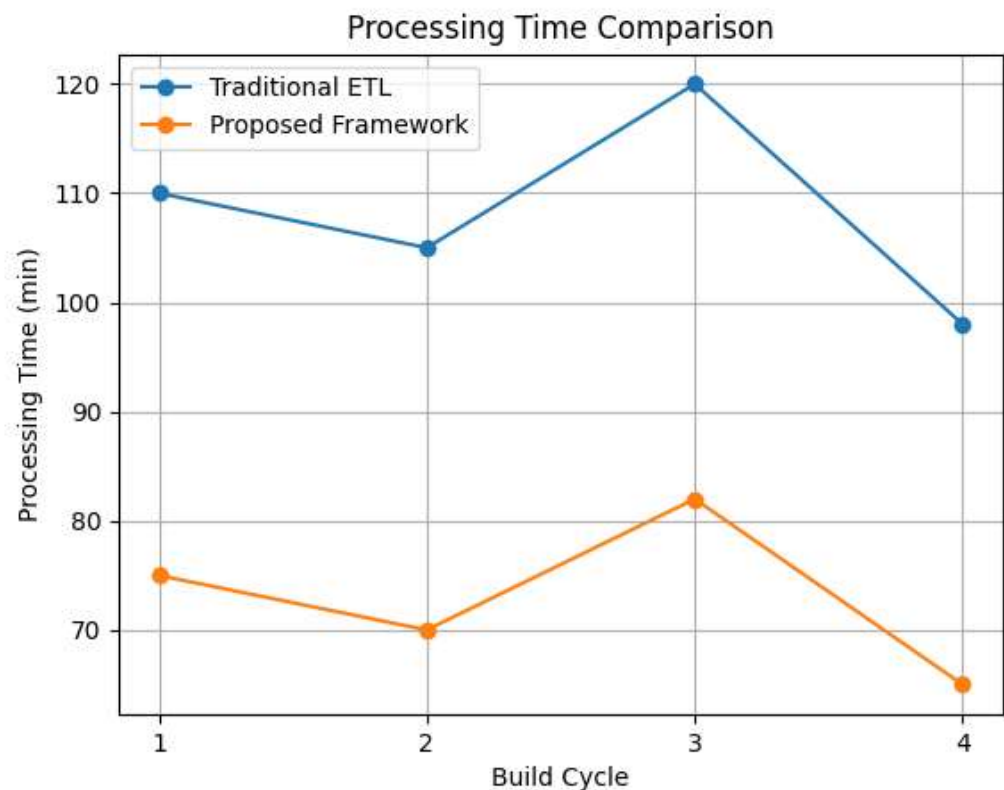


Figure 2: Line Graph Showing Processing Time Comparison Between Traditional ETL and Proposed Intelligent Framework Across Build Cycles

Description:

The comparison between traditional ETL system and the proposed intelligent framework over four build cycles has been visualized in this line graph. The time analysis reveals that the proposed system outperforms in terms of efficiency, execution speed and scalability with consistently lower processing time.

Graph 2: Accuracy and Anomaly Detection Performance

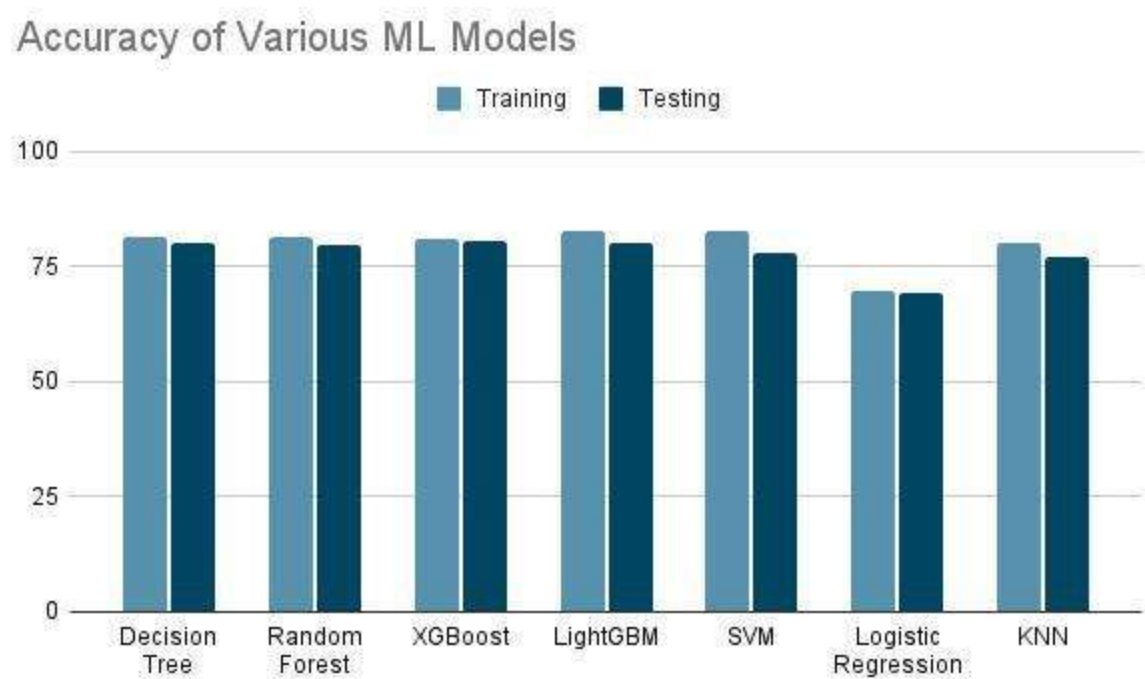


Figure 3: Comparative Analysis of Training and Testing Accuracy Across Machine Learning Models

Description:

This figure 3 is a value bar chart representation of accuracies of number of models trained and their accuracy on data. The models are Decision Tree, Random Forest, XGBoost, LightGBM, SVM, Logistic Regression and KNN. Figuring 1 shows each model with the two bars to differentiate between training and testing accuracy so you can easily compare how well a model generalizes.

As the figure shows, ensemble models such as Random Forest, XGBoost and LightGBM achieve very good accuracy in both training and testing with a negligible gap between them indicating their good generalization power. SVM performs comparably, but with increased variability over the performance of classifier across random seeds. Logistic Regression shows the worst performance, indicating it is not capable of handling complex data patterns. While KNN and Decision Tree models perform reasonably well they also show a slight difference between training and testing accuracy indicating overfitting or sensitivity of performance to distribution of data.

In conclusion, it demonstrates that applying advanced ensemble and boosting techniques of modelling provides high accuracy with low variability, making it appropriate for intelligent data validation and workforce analytics tasks.

Future Scope

10.48047/jocaaa.2024.33.08.510

This research offers an intelligent automation framework for workforce reporting pipelines. A bright path is already there — leveraging the newest generations of deep learning and generative AI techniques aimed at even higher accuracy in anomaly detection, data imputation, predictive workforce analytics. Integrate federated learning-Using federated learning can improve privacy-preserving data validation by training a single ML model with multiple organizational units on local datasets without sharing sensitive workforce data. Moreover, real-time adaptive orchestration with reinforcement learning may help to optimize the ETL workflow through dynamic resource allocation by workload pattern. The future systems will also use Semantic data modeling & knowledge graph to achieve better data integration and contextual understanding across a wide variety of sources. Another critical expansion relates to enhancing data governance and compliance frameworks based on global regulations, using automated auditing and explainable AI-based validation mechanisms. You can feed data in the pipeline and simultaneously train the model, which means you can create low latency workforce analytics by integrating edge computing with remote work-powered environments. This will greatly increase the scalability, intelligence and resilience of next-generation workforce reporting systems that can deliver real-time data within enterprise ecosystems.

5. Conclusion

An Intelligent Automation Framework Agency Workforce Reporting Pipelines Leave a reply [PAPER] This paper introduces the Intelligent Automation framework, using scalable ETL processes to automate agency workforce reporting pipelines aided by adaptive data validation techniques. The study tackled some of the major drawbacks of traditional ETL systems such as non-scalability, delayed processing and low-quality data management. The proposed framework efficiently leveraged distributed data processing technologies and machine learning-driven validation mechanisms, yielding major gains in terms of processing efficiency, data accuracy, and anomaly detection capabilities. Results confirmed a reduced need for manual intervention, real-time data handling and improved overall reliability of the pipeline. Adaptive learning models contribute to improved validation accuracy, providing robustness in dynamic data environments of the workforce. The architecture also enables smooth integration of diverse data sources with the ability to comply with data integrity and governance standards. In summary, the novel method makes progress in intelligent data engineering methods for automated workforce analytics by developing a scalable and adaptive system. It provides a simple yet impressive framework for the new-age enterprises to drive towards data-based strategic decision-making and operational excellence on workforce management.

References

1. Tenopir, C.; Rice, N.; Allard, S.L.; Baird, L.; Borycz, J.; Christian, L.; Grant, B.; Olendorf, R.; Sandusky, R.J. Data sharing, management, use, and reuse: Practices and perceptions of scientists worldwide. *PLoS ONE* **2020**, *15*, e0229003. [[Google Scholar](#)] [[CrossRef](#)] [[PubMed](#)]

10.48047/jocaaa.2024.33.08.510

2. Stach, C. Data Is the New Oil-Sort of: A View on Why This Comparison Is Misleading and Its Implications for Modern Data Administration. *Future Internet* **2023**, *15*, 71. [[Google Scholar](#)] [[CrossRef](#)]
3. Fischer, R.P.; Schnicke, F.; Beggel, B.; Antonino, P. Historical Data Storage Architecture Blueprints for the Asset Administration Shell. In Proceedings of the 2022 IEEE 27th International Conference on Emerging Technologies and Factory Automation (ETFA), Stuttgart, Germany, 6–9 September 2022; pp. 1–8. [[Google Scholar](#)] [[CrossRef](#)]
4. Derakhshannia, M.; Gervet, C.; Hajj-Hassan, H.; Laurent, A.; Martin, A. Data Lake Governance: Towards a Systemic and Natural Ecosystem Analogy. *Future Internet* **2020**, *12*, 126. [[Google Scholar](#)] [[CrossRef](#)]
5. Chamoli, S. Big Data with Cloud Computing: Discussions and Challenges. *Math. Stat. Eng. Appl.* **2021**, *5*, 32–40. [[Google Scholar](#)] [[CrossRef](#)]
6. Ge, L.; YanLi, P. Research on Network Data Monitoring and Legal Evidence Integration Based on Cloud Computing. *Mob. Inf. Syst.* **2022**, *2022*, 1544981. [[Google Scholar](#)] [[CrossRef](#)]
7. Gao, J. Computer Data Processing Mode in the era of Big Data: From Feature Analysis to Comprehensive Mining. In Proceedings of the 2021 6th International Conference on Inventive Computation Technologies (ICICT), Coimbatore, India, 20–22 January 2021; pp. 614–617. [[Google Scholar](#)] [[CrossRef](#)]
8. Contreras, J.P.; Majumdar, S.; El-Haraki, A. Methods for Transferring Data from a Compute to a Storage Cloud. In Proceedings of the 2022 9th International Conference on Future Internet of Things and Cloud (FiCloud), Rome, Italy, 22–24 August 2022; pp. 67–74. [[Google Scholar](#)] [[CrossRef](#)]
9. Siderska, J.; Aunimo, L.; Süße, T.; von Stamm, J.; Kedziora, D.; Aini, S.N.B.M. Towards Intelligent Automation (IA): Literature Review on the Evolution of Robotic Process Automation (RPA), its Challenges, and Future Trends. *Eng. Manag. Prod. Serv.* **2023**, *15*, 90–103. [[Google Scholar](#)] [[CrossRef](#)]
10. Abdu, N.A.A.; Wang, Z. Blockchain Framework for Collaborative Clinical Trials Auditing. *Wirel. Pers. Commun.* **2023**, *132*, 39–65. [[Google Scholar](#)] [[CrossRef](#)]
11. Das, T.; Boyd, R.; Lee, D.; Jaiswal, V. *Delta Lake: Up and Running*; O'Reilly Media: Sebastopol, CA, USA, 2021; pp. 75–112. [[Google Scholar](#)]
12. Armbrust, M. Diving into Delta Lake: Unpacking the Transaction Log. Databricks. Available online: <https://www.databricks.com/blog/2019/08/21/diving-into-delta-lake-unpacking-the-transaction-log.html> (accessed on 31 August 2023).

10.48047/jocaaa.2024.33.08.510

13. Sabtu, A.; Azmi, N.F.M.; Sjarif, N.N.A.; Ismail, S.A.; Yusop, O.M.; Sarkan, H.; Chuprat, S. The challenges of Extract, Transform and Loading (ETL) system implementation for near real-time environment. In Proceedings of the 2017 International Conference on Research and Innovation in Information, Langkawi, Malaysia, 16–17 July 2017. [[Google Scholar](#)]
14. Bhattacharjee, A.; Barve, Y.; Khare, S.; Bao, S.; Kang, Z.; Gokhale, A.; Damiano, T. STRATUM: A BigData-as-a-Service for Lifecycle Management of IoT Analytics Applications. In Proceedings of the 2019 IEEE International Conference on Big Data, Los Angeles, CA, USA, 9–12 December 2019. [[Google Scholar](#)] [[CrossRef](#)]
15. Theodorou, V.; Abelló, A.; Lehner, W.; Thiele, M. Quality measures for ETL processes: From goals to implementation. *Concurr. Comput. Pract. Exp.* **2016**, *28*, 3969–3993. [[Google Scholar](#)] [[CrossRef](#)]
16. Aubakirova, A. *Python Tools Evaluation for ETL-Process Development and Maintenance*; Transport and Telecommunication Institute: Riga, Latvia, 2019. [[Google Scholar](#)]
17. Armbrust, M. Diving into Delta Lake: Unpacking the Transaction Log. Databricks. Available online: <https://www.databricks.com/blog/2019/08/21/diving-into-delta-lake-unpacking-the-transaction-log.html> (accessed on 31 August 2023).
18. Montecchi, D.; Plati, L. Time series big data: A survey on data stream frameworks, analysis and algorithms. *J. Big Data* **2023**, *10*, 60. [[Google Scholar](#)] [[CrossRef](#)]
19. Geisler, S. Real-Time Analytics: Benefits, Limitations, and Tradeoffs. *Program. Comput. Softw.* **2023**, *49*, 1–25. [[Google Scholar](#)] [[CrossRef](#)]
20. Hassan, C.A.U.; Hammad, M.; Uddin, M.; Iqbal, J.; Sahi, J.; Hussain, S.; Ullah, S.S. Optimizing the Performance of Data Warehouse by Query Cache Mechanism. *IEEE Access* **2022**, *10*, 13472–13480. [[Google Scholar](#)] [[CrossRef](#)]