

International Journal of Data Engineering and Intelligent Computing

Vol. 1, No. 2, 2018

Doi: [10.67228/30715717/IJDEIC-2018PII2F6A](https://doi.org/10.67228/30715717/IJDEIC-2018PII2F6A)

PP. 01-12

Original Article

Automated Data Quality Monitoring in Streaming Data Transformations

Kevin Taylor¹, Dr. Ali Reza²

¹ Senior Project Manager, Prime Innovations, USA.

² Professor, University of Tehran, Iran.

Received: 06-02-2018

Revised: 06-12-2018

Accepted: 06-24-2018

Published: 07-06-2018

ABSTRACT

As real-time data becomes integral to decision-making, ensuring the quality of streaming data is critical to maintaining system reliability and business value. This paper explores automated data quality monitoring techniques specifically tailored for streaming data transformations. It investigates common data quality issues such as latency, duplication, schema drift, null values, and data inconsistency in real-time pipelines. The study presents a framework for integrating automated monitoring tools within streaming architectures using technologies like Apache Kafka, Apache Flink, and Spark Streaming. Machine learning-based anomaly detection models and rule-based engines are examined for detecting quality issues in-flight. Through experimental evaluation, the paper demonstrates how automated monitoring improves data observability, reduces downtime, and supports proactive data governance. The findings provide a roadmap for organizations seeking to implement scalable, real-time data quality solutions within their streaming ecosystems.

KEYWORDS

Streaming Data, Data Quality Monitoring, Real-Time Data Pipelines, Apache Kafka, Anomaly Detection, Schema Drift, Spark Streaming, Data Observability, Automated Validation, Streaming Analytics.

1. INTRODUCTION

1.1. Background and Motivation

In today's data-driven world, organizations increasingly rely on real-time data to inform critical business decisions, power operational dashboards, and drive machine learning pipelines. This surge in demand has popularized the use of streaming data platforms such as Apache Kafka, Apache Flink, and Spark Streaming, which enable continuous ingestion and transformation of data as it flows through the system. However, the velocity, variety, and volume of streaming data introduce significant challenges in ensuring its quality. Unlike traditional batch processing systems, where data can be inspected and cleaned before use, streaming systems process data in motion, leaving limited opportunity for manual validation or correction. As a result, even minor data quality issues such as null values, duplicate records, or schema mismatches can cascade into serious operational and analytical problems. This evolving context underscores the need for automated, real-time data quality monitoring to safeguard the integrity and usability of streaming data.

1.2. Importance of Data Quality in Streaming Systems

High-quality data is the foundation of reliable analytics and intelligent decision-making. In streaming systems, where data is consumed and acted upon in real-time, the importance of data quality becomes even more pronounced. Poor data quality in such systems can lead to inaccurate insights, flawed predictions, or even critical failures in applications such as fraud detection, financial trading, or predictive maintenance. Furthermore, streaming systems are often used in mission-critical environments where latency, throughput, and consistency must be guaranteed. Any deviation in expected data patterns such as missing timestamps, unexpected nulls, or corrupt records can impact downstream consumers instantaneously. Therefore, implementing robust data quality controls in streaming pipelines is not just beneficial but essential to ensure trustworthy, actionable outputs and to maintain the overall health of data-driven systems.

1.3. Challenges in Monitoring Real-Time Data

Monitoring data quality in streaming environments poses unique challenges that differ significantly from traditional batch systems. First, the continuous and high-throughput nature of streaming data means that any monitoring or validation must occur with minimal latency to avoid bottlenecks. Second, streaming pipelines often involve multiple stages of transformation and enrichment, each introducing potential points of failure or inconsistency. Moreover, streaming data is often semi-structured or unstructured, making it more susceptible to schema drift or unexpected changes in format. The lack of a static data snapshot further complicates the process, as it is difficult to apply retrospective quality checks. Additionally, streaming systems typically integrate with heterogeneous data sources, leading to variability in data standards and expectations. All these factors make real-time data quality monitoring a complex but necessary endeavor requiring automated, intelligent, and scalable solutions.

1.4. Objectives and Contributions

This paper aims to develop a comprehensive understanding of how automated systems can be employed to monitor and ensure data quality in streaming data transformations. The primary

objectives are to: (i) identify the most common data quality issues in streaming pipelines, (ii) evaluate current tools and techniques used for monitoring real-time data, and (iii) propose a scalable architecture that integrates automated rule-based and machine learning methods for continuous quality assessment. The paper's contributions include a detailed taxonomy of data quality dimensions in streaming environments, a comparative analysis of existing monitoring frameworks, and a novel approach for combining real-time alerts with historical analytics to detect and resolve data anomalies. By addressing both technical and operational aspects, the study provides a blueprint for organizations seeking to build resilient and trustworthy streaming data systems.

2. LITERATURE REVIEW

2.1. Traditional vs. Streaming Data Quality Approaches

In traditional data processing systems, data quality checks are typically performed in a batch-oriented manner. These systems rely on stable datasets that are processed in discrete intervals, allowing time for extensive validation, cleansing, and enrichment before the data is stored or analyzed. Data quality frameworks such as DataCleaner and Talend Data Quality operate effectively in such environments. However, with the advent of streaming data, these batch-based methods fall short. Streaming data pipelines require quality checks to be embedded within the transformation flow and executed in near real-time, often under strict latency constraints. Traditional approaches also lack the agility to adapt to rapidly changing schemas or high-frequency data updates common in streaming scenarios. As a result, there is a growing need to shift from static, retrospective data validation to dynamic, in-motion data quality monitoring that can react to anomalies as they occur.

2.2. Existing Tools and Frameworks (e.g., Deequ, Great Expectations, StreamAlert)

Several tools have emerged in recent years to address the growing need for data quality monitoring in both batch and streaming contexts. Deequ, developed by Amazon, provides a declarative framework for defining and validating data quality constraints in large-scale data processing systems. Although initially designed for batch workloads, Deequ can be adapted for micro-batching in Spark Streaming. Great Expectations is another popular open-source tool that allows for the creation of "expectations" about data values, types, and distributions. While powerful for batch validation, it offers limited native support for real-time streaming applications. In contrast, tools like StreamAlert and Satori have been designed with streaming environments in mind, offering built-in support for continuous validation, alerting, and integration with real-time monitoring dashboards. Despite these advancements, most tools struggle to maintain scalability, adaptability, and low-latency validation across diverse streaming platforms. This signals a gap in the availability of robust end-to-end solutions tailored for real-time data quality assurance.

2.3. Gaps in Current Research

While there is substantial academic and industrial interest in data quality, much of the existing research focuses on static or batch-oriented systems. Few studies delve into the unique challenges posed by real-time data streams, particularly in large-scale distributed environments. Key gaps include the lack of scalable frameworks for schema drift detection in real time, insufficient methods for proactive anomaly detection using streaming machine learning models, and limited

support for auto-remediation of data quality issues without manual intervention. Moreover, existing literature often overlooks the integration of data quality monitoring with operational monitoring systems like Prometheus, Grafana, or cloud-native observability stacks. This disconnect hinders real-time visibility into quality metrics. Bridging these gaps requires a combination of algorithmic innovation, architectural design, and domain-specific implementation strategies tailored to the streaming context.

3. DATA QUALITY DIMENSIONS IN STREAMING CONTEXT

3.1. Accuracy and Consistency

Accuracy refers to the degree to which streaming data correctly represents the real-world phenomena it is intended to model. Inconsistent or inaccurate data can arise from malformed inputs, sensor errors, or transformation logic failures during ingestion. Consistency, on the other hand, denotes the uniform application of rules and formats across the data stream. For instance, if a timestamp field sometimes contains Unix epoch values and other times ISO 8601 strings, it undermines the consistency of the stream. In streaming environments, maintaining accuracy and consistency is especially challenging due to the distributed nature of data sources and the rapid pace of updates. Automated monitoring must therefore validate value ranges, data types, and referential integrity in real time to prevent erroneous data from corrupting downstream systems.

3.2. Completeness and Timeliness

Completeness measures whether all expected data points are present in the stream. Missing fields, incomplete records, or lost messages due to buffer overflows can compromise the completeness of a streaming pipeline. Timeliness refers to the prompt arrival of data relative to its event time. In real-time analytics, late-arriving or delayed data can result in skewed aggregations and incorrect conclusions. Ensuring completeness and timeliness involves tracking data arrival rates, comparing against expected volumes, and flagging delays that exceed acceptable thresholds. Tools like watermarking in Apache Flink and event-time windows in Spark Streaming can assist in managing out-of-order or delayed data, but proactive quality checks are still needed to monitor and resolve such issues on the fly.

3.3. Schema Validation and Drift Detection

In a streaming context, data schemas may evolve over time due to upstream changes, such as modifications in application logic or new IoT device firmware. Schema drift refers to these changes in structure like renamed fields, altered data types, or added/removed attributes that can break downstream transformations and cause data loss. Schema validation involves ensuring that each message conforms to the expected schema before further processing. Automated schema validation engines can compare incoming records against a predefined contract, while drift detection systems monitor the frequency and nature of schema changes. Solutions like Confluent Schema Registry help manage schema evolution in Kafka-based systems, but real-time validation and impact assessment remain key challenges that require further innovation.

3.4. Duplication and Anomaly Detection

Duplicate records in streaming data pipelines can arise from network retries, partition rebalancing, or poorly designed transformation logic. These duplicates, if undetected, can lead to inflated counts, misleading metrics, or inconsistent state updates. Anomaly detection, meanwhile, focuses on identifying data points that deviate from expected patterns such as sudden spikes in temperature readings, suspicious transaction amounts, or irregular data arrival intervals. Real-time anomaly detection often relies on lightweight statistical models, machine learning algorithms, or hybrid rule-based systems capable of operating with minimal latency. Implementing automated deduplication and anomaly detection is crucial for maintaining the credibility and operational reliability of streaming systems, especially in high-stakes domains like finance, healthcare, or cybersecurity.

4. PROPOSED ARCHITECTURE

4.1. System Overview

The proposed architecture for automated data quality monitoring in streaming environments is designed to operate in real time, integrating seamlessly with existing data pipelines without introducing significant latency. The architecture is modular and scalable, built on top of widely adopted open-source technologies that support high-throughput and low-latency data processing. At a high level, the system consists of three key layers: the data ingestion layer, the transformation engine, and the monitoring layer. Data flows from source systems through ingestion tools such as Apache Kafka or AWS Kinesis, where it is partitioned and buffered for downstream processing. The transformation layer, powered by stream processing engines like Apache Flink or Spark Structured Streaming, applies business logic and enriches the data. The monitoring layer, which sits alongside or within the transformation logic, inspects the data in motion to detect quality issues using both rule-based and machine learning techniques. This layered approach ensures that data quality checks can be distributed, fault-tolerant, and continuously aligned with evolving data structures and business requirements.

4.2. Components

4.2.1. Data Ingestion Layer (*Kafka, Kinesis*)

This layer is responsible for capturing raw data from source systems in real time and ensuring reliable delivery to processing components. Apache Kafka and AWS Kinesis are popular choices due to their durability, partitioning, and replay capabilities. In the context of data quality, the ingestion layer provides the first opportunity to capture metadata such as timestamps, message sizes, and schema identifiers, which are essential for downstream validation. Kafka's integration with Schema Registry also allows for schema enforcement at the topic level, reducing the likelihood of malformed records entering the pipeline.

4.2.2. Transformation Engine (*Flink, Spark*)

Once data is ingested, it is handed off to a stream processing engine for transformation. Apache Flink and Spark Streaming offer robust APIs for processing streaming data in event time or processing time, applying windowed aggregations, joins, filters, and complex event processing logic.

These engines also allow for embedding validation logic within the processing flow, enabling the system to inspect each record or batch for conformance to defined quality rules. For instance, Flink's DataStream API can be extended with side outputs to flag and isolate invalid records without interrupting the main dataflow, thus supporting real-time remediation or routing to quarantine queues.

4.2.3. Monitoring Layer (Custom Scripts, ML Models, and Rule Engines)

The monitoring layer is the heart of the architecture, responsible for performing automated data quality checks. It comprises custom scripts, rule engines (such as Deequ or custom rule DSLs), and machine learning models that continuously analyze data for anomalies, schema deviations, null values, duplicates, and more. Rule engines can validate fields against predefined constraints (e.g., numeric ranges, allowed values), while machine learning models like Isolation Forest or Autoencoders detect outliers or unusual patterns in the data stream. The monitoring layer is tightly integrated with observability tools and metrics collectors to ensure full transparency and traceability of quality issues as they arise.

4.3. Integration with CI/CD Pipelines and Alerting Systems

For the architecture to be production-ready and maintainable, it must support integration with continuous integration and deployment (CI/CD) pipelines. Data quality rules, transformation logic, and monitoring configurations can be version-controlled and tested just like application code, enabling automated deployments and rollbacks. Infrastructure-as-code tools like Terraform or Helm can be used to manage the deployment of streaming components. Additionally, the monitoring layer should emit metrics and alerts to external systems such as Prometheus, Grafana, or cloud-native services like AWS CloudWatch or Google Cloud Operations. These integrations allow teams to visualize real-time quality metrics, receive alerts for critical violations, and respond quickly to data issues before they propagate to downstream systems.

5. METHODOLOGIES FOR AUTOMATED MONITORING

5.1. Rule-Based Quality Checks (Thresholds, Patterns)

Rule-based quality checks represent the most interpretable and deterministic form of data validation. These checks are defined based on domain knowledge, business requirements, and data contracts. Examples include verifying that a field is non-null, that a numeric value falls within a specific range, or that a categorical field only contains allowed values. Rules can also include more complex patterns, such as verifying that event timestamps increase monotonically or that IDs are unique within a given window. In streaming systems, these rules are applied inline during processing or as part of dedicated validation branches within the dataflow. The rules are typically written as declarative configurations or embedded as logic in transformation scripts. Rule violations can trigger alerts, be routed to dead-letter queues, or be logged for further analysis. Rule engines like Great Expectations or Deequ can be adapted for micro-batch or streaming use cases, enabling users to maintain a living set of expectations that evolve alongside the data.

5.2. ML-Based Anomaly Detection (Isolation Forest, Autoencoders)

Machine learning techniques are increasingly used to detect subtle or previously unseen data quality issues that are not captured by static rules. In streaming contexts, models like Isolation Forest, One-Class SVMs, or deep learning-based Autoencoders can be trained on historical data to learn the “normal” distribution and then applied in real time to flag deviations. Isolation Forest, for instance, isolates anomalies by randomly partitioning the data space, assigning higher anomaly scores to outliers. Autoencoders learn to compress and reconstruct normal data and are effective in identifying anomalous records based on reconstruction error. These models can run as lightweight inference services that consume streaming data, producing real-time flags or scores for integration with alerting systems. While ML-based monitoring introduces additional complexity and compute cost, it greatly enhances the system’s ability to detect nuanced or evolving data issues without prior definition.

5.3. Real-Time Dashboards and Alerting Mechanisms

To operationalize data quality monitoring, it is essential to visualize the state of the system and surface alerts to stakeholders. Real-time dashboards built with tools like Grafana, Kibana, or cloud-native visualization platforms allow users to observe metrics such as rule failure rates, schema changes, or anomaly scores over time. Dashboards can be configured with thresholds to automatically trigger alerts via email, Slack, PagerDuty, or other incident management tools. For instance, a spike in null value rates or a sudden drop in data volume can immediately raise a red flag. These dashboards support filtering, drill-downs, and historical comparisons, enabling both reactive and proactive monitoring. Moreover, integration with observability platforms ensures that quality monitoring becomes a first-class citizen in the broader operational health of data pipelines.

5.4. Logging and Audit Trail Generation

An often overlooked but critical aspect of data quality monitoring is maintaining a comprehensive audit trail of all validation events, rule violations, schema changes, and remediation actions. In a streaming context, this involves writing structured logs to centralized systems like Elasticsearch, AWS CloudWatch, or Kafka log topics dedicated to audit data. Each log entry can include metadata such as the record ID, the violated rule or anomaly type, timestamp, and action taken. These logs not only support root-cause analysis during incidents but also provide an essential foundation for compliance reporting and post-mortem reviews. Audit trails can also be analyzed retrospectively to fine-tune monitoring strategies, evaluate false positive rates, and identify recurring patterns of data degradation.

6. CASE STUDY/EXPERIMENTAL EVALUATION

6.1. Dataset and Streaming Setup

To evaluate the effectiveness of the proposed monitoring system, a case study was conducted using a simulated financial transaction dataset streamed in real time. The dataset includes user transactions with attributes such as transaction ID, amount, timestamp, location, and payment method. The streaming setup consists of Apache Kafka as the ingestion layer, Apache Flink for transformation and processing, and a custom monitoring layer incorporating both rule-based checks

and ML-based anomaly detection. The environment was deployed using Docker containers orchestrated via Kubernetes, enabling reproducibility and scalability of experiments. Data was generated using a synthetic data generator simulating real-world patterns, including periodic anomalies and schema changes to evaluate robustness.

6.2. Implementation of Monitoring Rules and Models

Rule-based checks were implemented using a custom validation library integrated into the Flink pipeline, covering constraints like non-null values, valid currency codes, and positive transaction amounts. Simultaneously, an Isolation Forest model was trained on a baseline of “normal” transaction data and deployed as a microservice for real-time scoring. Each incoming record was evaluated by both the rule engine and the anomaly detection model, and results were routed to Kafka topics dedicated to alerts and logs. The system also employed a schema registry to detect field additions, deletions, or type changes in the transaction schema, with validation logic triggering alerts upon deviations.

6.3. Evaluation Metrics (Precision, Latency, Accuracy)

The monitoring system was evaluated using standard performance metrics. Precision measured the proportion of correctly identified quality issues among all flagged records, while accuracy evaluated overall detection performance across the entire stream. Latency was measured as the time between data ingestion and quality alert generation. The rule-based system achieved high precision (>95%) for well-defined violations, while the ML model demonstrated improved sensitivity in detecting subtle outliers that the rules missed. The combined architecture introduced an average processing latency of less than 150 milliseconds, well within acceptable limits for real-time applications.

6.4. Results and Insights

The case study demonstrated the value of combining rule-based and machine learning approaches for comprehensive data quality monitoring in streaming pipelines. Rule-based checks were effective for enforcing static data contracts, while the ML component enhanced adaptability to new or unexpected data patterns. The monitoring layer’s integration with dashboards and alerting systems provided visibility into real-time data health and facilitated rapid troubleshooting. Additionally, the experiment highlighted the importance of schema validation and audit logging in maintaining operational confidence. Overall, the results confirmed the feasibility and utility of automated, real-time data quality monitoring as an integral part of modern streaming data architectures.

7. DISCUSSION

7.1. Benefits of Automation in Streaming Quality

Automating data quality monitoring in streaming environments delivers significant benefits over manual or ad hoc approaches. One of the most important advantages is real-time responsiveness automated systems can detect and flag issues within milliseconds of data arrival, allowing for immediate intervention. This proactive capability helps prevent the propagation of bad data into

downstream systems such as dashboards, machine learning models, or data warehouses. Furthermore, automation ensures consistency in the application of data quality rules, removing the variability associated with manual reviews. It also allows for continuous coverage at scale, handling large volumes of high-velocity data that would be impossible to validate manually. The integration of machine learning into automated systems enhances adaptability by identifying previously unseen data anomalies, thus increasing resilience against evolving data patterns and schema changes. Overall, automation enables organizations to maintain trust in their real-time data pipelines while reducing operational overhead and improving decision-making accuracy.

7.2. Limitations and Challenges

Despite its advantages, automated data quality monitoring in streaming systems faces several challenges. First, there is a trade-off between validation accuracy and system latency excessively complex validation logic or heavy model inference can slow down processing and affect real-time performance. Second, designing effective quality rules requires domain expertise, and over-reliance on predefined rules may lead to blind spots or brittle systems. ML-based detection mechanisms, while powerful, are not immune to false positives and require continuous retraining and validation to remain effective. Another limitation is the integration complexity, especially in heterogeneous environments with diverse data sources, transformation engines, and monitoring platforms. Ensuring that monitoring tools are interoperable and scalable across cloud and on-premise systems adds another layer of complexity. Additionally, schema evolution and backward compatibility in streaming systems present persistent challenges for rule enforcement and anomaly detection. Finally, data privacy and compliance concerns may arise when storing and analyzing sensitive data for quality assessment, especially when audit logs and raw payloads are retained for traceability.

7.3. Comparison with Batch-Based Quality Monitoring

Compared to traditional batch-based monitoring systems, streaming quality monitoring introduces a paradigm shift in how data quality is perceived and managed. Batch systems benefit from access to complete datasets, allowing for comprehensive validation, retrospective corrections, and complex dependency analysis. However, they often operate with significant delays and are unsuitable for real-time applications. In contrast, streaming systems prioritize immediacy and continuity, requiring validation to occur on-the-fly without access to the full dataset. While this improves timeliness, it limits the complexity and depth of validations that can be performed. Streaming systems also face unique issues like out-of-order events, data duplication due to retries, and time-window-based aggregation errors, which are less common in batch workflows. Despite these differences, the growing convergence between batch and stream processing exemplified by technologies like Apache Beam indicates a trend toward unified quality monitoring frameworks that leverage the strengths of both paradigms.

8. FUTURE WORK

8.1. Adaptive Learning for Quality Rules

One promising direction for future research is the development of adaptive learning mechanisms for data quality rules. Currently, most rule-based systems rely on static configurations

that must be manually defined and updated. However, with the growing complexity and dynamism of streaming data, a more flexible approach is needed. Adaptive learning involves using feedback loops, reinforcement learning, or pattern mining to automatically refine or generate quality rules based on historical validation outcomes and evolving data characteristics. Such systems could learn from labeled anomalies, analyst interventions, or user feedback to improve detection precision and reduce false positives. Integrating adaptive learning with streaming engines would allow for self-healing pipelines that evolve in real time, improving long-term resilience and minimizing the need for human intervention.

8.2. Integration with Data Lineage Tools

Another crucial area for advancement is the integration of streaming quality monitoring with data lineage and metadata management tools. Data lineage provides a visual and programmatic trace of data from its source through all stages of transformation to its final destination. Coupling lineage with quality monitoring enables organizations to understand not just when a data issue occurred, but where and why. This end-to-end traceability is essential for root cause analysis, regulatory compliance, and impact assessment. Emerging standards like OpenLineage and tools such as Apache Atlas and DataHub provide the infrastructure to support this integration. Future systems should enable real-time lineage updates, correlating quality alerts with specific transformation steps or schema changes, thereby accelerating debugging and governance workflows.

8.3. Edge-to-Cloud Streaming Quality Monitoring

As more data is generated at the edge from IoT devices, mobile sensors, or embedded systems there is an increasing need to extend quality monitoring capabilities beyond the cloud. Edge-to-cloud streaming introduces unique challenges such as intermittent connectivity, limited compute resources, and varied data standards. Future work should explore lightweight, decentralized quality checks that can be deployed directly on edge devices or gateways, with summary statistics and violations transmitted to the cloud for aggregation and alerting. This hybrid model would enable early detection of data issues closer to the source, reducing noise in the central pipeline and allowing for localized remediation. Additionally, adaptive sync mechanisms could ensure consistency between edge-validated data and cloud-based monitoring systems, forming a cohesive quality fabric across the data lifecycle.

9. CONCLUSION

9.1. Summary of Findings

This paper has examined the architecture, methodologies, and practical considerations of implementing automated data quality monitoring in streaming data transformations. By analyzing common data quality dimensions such as accuracy, timeliness, schema validity, and anomaly detection the study highlighted the complexities of maintaining data trustworthiness in high-velocity environments. It proposed a layered system architecture integrating Kafka/Flink with rule engines and ML models, capable of executing inline validations and delivering real-time alerts. A case study using financial transactions demonstrated the viability of the system, with high precision, low

latency, and improved data observability. These findings establish that with the right design and tooling, automated monitoring can significantly enhance data reliability in streaming systems.

9.2. Implications for Industry and Research

For industry practitioners, the findings underscore the necessity of embedding automated quality checks directly within streaming pipelines to mitigate the risks of real-time data corruption. The combination of rules and machine learning models offers a flexible and scalable approach suitable for various domains, including finance, e-commerce, healthcare, and manufacturing. For researchers, the study opens new directions for exploring self-learning quality systems, lineage-aware monitoring, and the challenges of extending these concepts to federated or edge environments. There is also a need for more standardized benchmarking datasets and evaluation frameworks to assess the performance of quality monitoring tools under different operational constraints.

9.3. Final Remarks

In conclusion, as organizations transition from batch-centric architectures to real-time, event-driven systems, data quality monitoring must evolve accordingly. Automation is not just a convenience it is a necessity for maintaining trust and compliance in modern data ecosystems. While challenges remain, the convergence of streaming analytics, machine learning, and observability tooling presents a powerful opportunity to build intelligent, self-aware data pipelines that monitor, adapt, and improve over time. Future work should aim to democratize access to these systems and make them more accessible, interpretable, and resilient across diverse data landscapes.

REFERENCES

- [1] Wang, R. Y., & Strong, D. M. (1996). Beyond accuracy: What data quality means to data consumers. *Journal of Management Information Systems*, 12(4), 5–33. <https://doi.org/10.1080/07421222.1996.11518099>
- [2] Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., Chaudhary, V., Young, M., Crespo, J. F., & Dennison, D. (2015). Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems*, 28, 2503–2511.
- [3] Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM Computing Surveys*, 41(3), Article 15. <https://doi.org/10.1145/1541880.1541882>
- [4] Zaharia, M., Das, T., Li, H., Hunter, T., Shenker, S., & Stoica, I. (2013). Discretized streams: Fault-tolerant streaming computation at scale. In *Proceedings of the Twenty-Fourth ACM Symposium on Operating Systems Principles* (pp. 423–438). <https://doi.org/10.1145/2517349.2522737>
- [5] Batini, C., & Scannapieco, M. (2016). *Data and information quality: Dimensions, principles and techniques*. Springer. <https://doi.org/10.1007/978-3-319-24106-7>
- [6] Kimball, R., & Ross, M. (2013). *The data warehouse toolkit: The definitive guide to dimensional modeling* (3rd ed.). Wiley.
- [7] Cugola, G., & Margara, A. (2012). Processing flows of information: From data stream to complex event processing. *ACM Computing Surveys*, 44(3), 1–62. <https://doi.org/10.1145/2187671.2187677>
- [8] Menascé, D. A., & Almeida, V. A. F. (2002). *Capacity planning for Web services: Metrics, models, and methods*. Prentice Hall.
- [9] Aggarwal, C. C. (2013). *Outlier analysis* (2nd ed.). Springer. <https://doi.org/10.1007/978-1-4614-6396-2>
- [10] English, L. P. (2009). *Information quality applied: Best practices for improving business information, processes and systems*. Wiley.
- [11] Akidau, T., Bradshaw, R., Chambers, C., Chernyak, S., Fernández-Moctezuma, R. J., Lax, R., McVeety, S., Mills, D., Perry, F., Schmidt, E., & Whittle, S. (2015). The dataflow model: A practical approach to balancing correctness, latency,

- and cost in massive-scale, unbounded, out-of-order data processing. Proceedings of the VLDB Endowment, 8(12), 1792–1803. <https://doi.org/10.14778/2824032.2824076>
- [12] Kreps, J., Narkhede, N., & Rao, J. (2011). Kafka: A distributed messaging system for log processing. Proceedings of the NetDB Workshop, 1–7.
- [13] Ahmed, M., Mahmood, A. N., & Hu, J. (2016). A survey of network anomaly detection techniques. Journal of Network and Computer Applications, 60, 19–31. <https://doi.org/10.1016/j.jnca.2015.11.016>
- [14] Dunning, T., & Friedman, E. (2014). Time series databases: New ways to store and access data. O'Reilly Media.
- [15] Dean, J., & Ghemawat, S. (2008). MapReduce: Simplified data processing on large clusters. Communications of the ACM, 51(1), 107–113. <https://doi.org/10.1145/1327452.1327492>