



SRE-Driven Reliability Framework for Enterprise Data Platforms

Praveen Kumar Aaviti

Analytics Solutions Manager / Data Owner, Leading US Bank, United States

Publication History: Received: 01.07.2026; Revised: 27.07.2026; Accepted: 01.08. 2026; Published: 12.08.2026.

ABSTRACT: Business data platforms, which exist today, can enable major business processes through the use of massive data streams, analytics platforms, dashboards, reports and products derived through big data. The growing amounts of data, distributed architecture and even intricate dependencies have however rendered it difficult to ensure that there is a consistent reliability and operational stability. This work suggests an SRE-based Reliability Framework of Enterprise Data Platforms, which would use the concepts of Site Reliability Engineering (SRE) to data engineering and analytics experience. Data reliability is operationally represented and numerically measured according to the proposed framework with the help of Service Level Indicators (SLIs), Service Level Objectives (SLOs) and error budgets. It evaluates the key metrics that include availability of pipelines, freshness of the data, quality of the data, processing latency, availability of the dashboards, report delivery and endurance of the analytics service. It is also seen within the framework that there is also a repetitive reliability lifecycle that includes reliability measurement, observability, anomaly detection, incident management, automated remediation and post incident learning. Under an error-budget mechanism, the organization is able to trade features development and stability in the operation by balancing out the changes as reliability goals are met each time. Dependency-aware monitoring and reliability scoring between interconnected data products and pipelines, is also introduced by the framework. The goals of this integration of engineering-data governance-operational analytics will be to decrease the number of incidents and enhance Mean Time to Detect (MTTD) and Mean Time to Recover (MTTR) and trust of enterprise data services. The schema offers a growing base towards attaining quantifiable, preventive and enduring information platform trustworthiness.

KEYWORDS: Site Reliability Engineering (SRE), Data Reliability, Service Level Objectives (SLOs), Error Budgets, Data Freshness, Pipeline Availability, Incident Management

I. INTRODUCTION

Business enterprises are increasingly turning to data platforms to help aid their decision making process in the areas of operations, business intelligence, customer services, financial analysis, regulatory reporting, artificial intelligence and strategic plans. The existing data platforms have now become many-layered systems with data lakes, lake house, streaming platform, cloud storage, data warehouse data transformation pipeline, APIs, dashboards, analysis applications and machine learning services becoming the transformations of the past centralized data warehouses. Although such technologies enable high scalability and flexibility, it is accompanied by a high degree of complexity of operation. A single upstream pipeline with some weaknesses, an unreasonably lengthy data ingestion rate, schema change, or data infrastructure, or external reliance can a trickle down into numerous data offerings by the business, eventually into business decisions. As such, reliability has now become a very important attribute of enterprise data platforms and not an issue pertaining to the infrastructure [1].

In the traditional process of data engineering development, the data quality verification and the fate of job completion is typically through pipeline building. Though these aspects are still very significant, they fail to offer a wholesome mechanism used to gauge the reliability as it feels to the end users and business applications. It will then feed (technically) old information to a dashboard using the example of a pipeline. Just as much a report-generation service will proceed to execute and generate part, or incorrect information as, due to a delay on an upstream set of data. Enterprise data reliability should then be measured on numerous aspects which encompasses: availability of pipelines; data freshness; data processing latency; data completeness; data quality; dashboard availability; report delivery and analytics service continuity [2][3].

Site Reliability Engineering (SRE) is one way of going about them. Theorized SRE was originally about the quality of reliability (Large-scale software service) management and its priorities are quantifiable reliability goals, automation,



observability, incident response, and continual improvement. Service Level Indicators (SLIs), Service Level Objectives (SLOs) and Service Level Agreement (SLAs), error budgets, monitoring, alerting and after incident analysis can be applied to data spaces of enterprises. By taking advantage of an adaptation of this nature, data pipelines and analytical services are the type of production services that organizations can consume with well-defined reliability expectations [4].

A SLI will provide a quantitative measure of a certain reliability property. Rate of successful pipeline executions, ratio of datasets that fulfil the freshness requirement, availability of dashboards, success rate of report delivery and latency are services that can be accessed as unlocks, in an enterprise data platform. Such measurements are then translated into quantifiable reliability objectives by SLOs. One definition of an example will be that an objective of 99.9 percent of the data pipelines of some significance are expected to finish in their expected processing time or 99 percent of business-important data sets are predicted to satisfy specific requirements of freshness. These goals establish a standard vocabulary where the reliability could be measured by data engineers and platform teams, application teams, and the business interests.

Error budget is another concept in which this approach is supported. Instead of absolute reliability Error budgets vary defined the maximum number of connection service degradation as associated with a given SLO. As an example, to the extent that there is the critical data pipeline, and the availability target of 99.9% in the scenario, the remaining fraction is the permissible reliability budget of failure or latitude. By employing the error budget in the short run, organizations are able to buy some time in order to focus on the development of reliability instead of new features. On the other hand, in cases where there is sufficient budget available new capabilities can be rolled out by the development teams with acceptable operational risk. It is in this process that a fair proportion between innovation and production is achieved.

Although the tenets of SRE apply, various problems have been associated with applying the general SRE models to the data platforms. Software services typically also take availability and response time into account, but data platforms must also consider such attributes as freshness, completeness, accuracy, lineage, consistency and successful propagation across dependent data products. A dashboard can be part of an infrastructure but the value of the dashboard could be lower since the data in the dashboard is not recent. Consequently, both the data reliability and service reliability need to be included in the effective framework of reliability.

To handle this requirement, this paper suggests a Reliability Framework of Enterprise Data Platforms, which is driven by SREs. The framework defines a single reliability model which ties together data pipelines, datasets, reports, dashboards, analytical services, and downstream data products with quantifiable reliability goals. It also deploys five main operation capabilities, either reliability measurement, continuous observability, automated remediation, intelligent incident detection and post-incident learning. Reliability measurements are elevated on various levels of data ecosystem and refined into valuable index of reliability. It keeps upstream failure/downstream degraded service dependency information to assist in identification of dependencies as well in increasing incident diagnoses and priority [5].

The combination of data governance and operational management and including reliability engineering as part of the proposed framework are also emphasized. The framework does not consider data quality, pipeline operations and business analytics as dissimilar functions, but it considers them as elements of a lifecycle of reliability. Where a reliability violation occurs, alerts brought up by monitoring mechanisms, incident-management procedures classify the incident and remediation procedures aid in making a quick recovery. The causes and post-incident lesson learned are subsequently analyzed and the organizations can be enhanced to monitor the rules, pipeline design, automation and SLOs.

One of the objectives of the proposed strategy is to move shift enterprise data operations towards reactive failure management in a proactive, quantitative manner that is, reliability engineering. The reliability metrics can help organizations view the frequency of failures in pipelines, freshness failure, and the downgrade of availability, and bottlenecks in an operations performance. Besides, other parameters like Mean Time to Detect (MTTD) and Mean Time to Recover (MTTR) that might be used to measure the efficiency of the monitoring and incident-response process, could be included into the reliability assessment.

It is a framework which is supposed to operate within large scale, nonhomogenous enterprise environments, with numerous data technologies and enterprise applications communicating in real time. Through the creation of a set of consistent reliability measures and targets, organizations will be able to measure the performance of a variety of data



products, make reliable investments, and make effective decisions on operational risk. The proposed SRE-oriented solution expands on the conventional details of data engineering, therefore, introducing a methodical process of metric continuously quantifying, executing and streamlining decision-worthiness of enterprise data services.

In the following sections, the structure of the framework, and an architectural building-block of the framework, reliability analysis, SLO and error-budget analysis, incident-management life-cycle of the framework and the approach of evaluation is analysed and presented. The goal of this paper is to explain how the adoption of the SRE principles and using the data engineering can offer the scalable backplane to make the production more stable, revitalize the research, offer availability to the pipeline, be resilient, and offer confidence in the enterprise analytics.

II. RELATED WORK

Enterprise data platform reliability has come to be a pertinent research topic as organizations have come to rely on data pipelines and analytics services, dashboards, and AI-enabled applications. The data observability, data quality, lineage, monitoring, governance, interoperability and responsible AI have been the focus of current research. Nonetheless, minimal efforts have been suggested to combine these attributes with the concepts of Site Reliability Engineering (SRE) like Service Level Indicators (SLIs), Service Level Objectives (SLOs), error budgets, incident management, Mean Time to Detect (MTTD) and Mean Time to Recover (MTTR). The proposed SRE-Driven Reliability Framework of Enterprise Data Platforms is based on the following studies.

Mekala [1] presented an observability model in AI-driven pipelines with the emphasis on a real-time view and debugging of multi-faceted AI processes. The paper underlines the importance of continually observing the behaviour of pipelines with the view to identifying any anomalies in the way they are operating and also facilitating the quick debugging. The ability to observe such can be very relevant to enterprise data platforms where a failure in upstream ingestion, transformation or processing can impact various downstream data products. This though is more related to monitoring, debugging, but not establishing specific, concrete reliability goals, nor ensuring reliability with SLOs and error budgets. The suggested framework builds on this view of observability by linking measurements of monitoring to formalized reliability targets and the operational response mechanisms.

Bangad et al. [2] offered a theoretical data quality monitoring framework of high-volume data settings that is AI-driven. They demonstrate the difficulties of size of data quality monitoring in their work and or they are focused on the automated responses to detecting quality issues. The quality of business data is a factor that should include data quality since whilst the infrastructure supporting the enterprise may be operating at optimum rate, inadequate data, or partial data can compromise the statistical results of analysis. The article offers an easy base to automated data-quality monitoring, yet the production service and data availability, incident management, error-budgets and recovery-goals are not completely met. The suggested framework therefore, includes data-quality monitoring as an element of a broadened example of reliability.

Patel, Gupta and Shaw [3] examined common data-quality checks of data-Mesh systems that employed AI. Their work is especially topical since the data mesh environments spread the data ownership among their areas, enhancing the significance of unified quality controls. Validation mechanisms should be consistent in order to enhance the reliability between data products, which are independently handled by data products. Nevertheless, functional reliability lifecycle is not defined but when quality checks have a defined standard. Mechanisms in enterprise platforms are needed to quantify availability, freshness, and latency, incident frequency, and recovery performance on top of data quality. The suggested framework helps to eliminate this flaws by expanding the standardized quality surveillance to establish the SLO-based quality management.

Another requirement that should be given serious consideration in the presence of the dependable enterprise data systems and that is the data lineage is listed below. Hummer, Green and Lee [4] have talked about data-intensive applications based on lineage techniques, by saying that one should be aware of how data sources, transformations, and consumers that absorb the data interact. The lineage information can be used to dramatically enhance incident investigation since engineers can know which datasets, reports, dashboards, or analytical services can be impacted by a failure upstream. The lineage can be considered as one of such dependency-awareness mechanisms that enable the impact analysis, the incident prioritization and the exploration of the root causes in the proposed model. As such, lineage does not just concern the governance but also the reliability of the operation.



Liu and Smith [5] subject AI to testing whereby it takes real-time scanning measurements to verify data integrity. Their experiment suggests identifying anomalies and data integrity problems in data environments via data environments diligence with resorting to AI-based solutions. Enterprise platforms are particularly important during real-time monitoring as any moment of failing to identify issues of data will lead to wrongful calculations of reports and business decisions. But, making checks on information integrity is merely one element in end-to-end reliability. The proposed framework adds integrity monitoring, the availability of pipeline, availability of data, processing latency, incident response, and availability of analytics to gain a broader perspective of the reliability.

The standard data-quality characteristics is suitable ISO / IEC 25012 standard data quality model [6]. Specifying systematic data-quality dimensions and outlining measurable, data-asset-assessment requirements can help. The proposed data service quality quality and its standardized characteristics can be used to enhance the SRE indicators in the proposed framework by the specificity of the technical availability of a data service and quality of data service provided by the specific data service. The difference between the two is huge since a pipeline that runs continuously might contradict business requirements when it provides incomplete, disjointed and stale information.

The Internet-of-Things systems are architectural with IEEE Std 2413-2022 [7] being applicable to the proposed work since enterprises data platforms are involved in the merging of heterogeneous devices, sensors, applications, and data sources. These distributed setups need a methodical architectural and functional strategy in dealing with reliance and information streams. Similarly, the IEEE P2302 standard [8] has to do with interoperability and portability of cloud computing. Its principles have been shown to apply to enterprise data platforms across a broad array of cloud environments as the reliability can be impaired by infrastructure heterogeneity, intercloud portable constraints and dependences. Such standards, however, are mostly concerned with architectural interoperability as opposed to operational SRE metrics.

The NIST Artificial Intelligence Risk Management Framework (AI RMF) [9] says it offers a systematic approach to identifying, assessing and managing AI system-related risks. Its areas of focus on the governance, measurement, and the ongoing risk management apply to the data reliability in businesses, especially when AI is applied to identify anomalies, as a proactive measure, or to make automatic corrections. The suggested framework may be added to this risk-focused view that would encourage indicators of operational reliability and SLO based controls to data services of the enterprise overall.

The interest of transparency of autonomous systems is the IEEE Std P7001-2022 [10]. It is reasonable to provide transparency in cases where automated monitoring and remediation models are deployed into enterprise environments as organizations in operation must have sufficient visibility of the causes of an automated action. The suggested framework will consider the observability and evidence of the incident, which will be part of traceable reliability operations to increase accountability in regards to automated decisions.

The opportunity of AI implementation in the quality monitoring analytics is the subject of one of the studies conducted by Shah, Gami, and Trehan [11]. Their effort enhances the contribution of smart approaches in identifying data-quality concerns and facilitate analytics dependability. But data quality does not relate to the holistic user experience of enterprise data services. The proposed framework, thus, has quality indicators, availability, freshness, latency, incident response and recovery metrics.

An organizational lens to evaluate the use and maturity of AI capabilities is a maturity model of AI capabilities that is offered by Gartner [12]. The development of organizations can be studied using it, as these have initially engaged in simple experimentation but have progressed through time to become operationally mature AI environments. Maturity evaluation however does not identify quantitative data pipeline and analytical service reliability goals.

EU AI act [13] and the General Data Protection Regulation (GDPR) [14] provide the regulatory terms of the systems conscious and reputable AI systems and privacy and security of personal information, respectively. These governance mechanisms suggest that the enterprise data platforms need to emphasise on governance, accountability, security and compliance as well as technical reliability. Regulatory compliance however, is not a direct source in provision of operation SLO management mechanisms and error-budget-based-decision-making mechanisms.



III. SRE-DRIVEN RELIABILITY FRAMEWORK FOR ENTERPRISE DATA PLATFORM

The proposed SRE-based Reliability Framework of Enterprise Data platforms leverages the concepts of Site Reliability Engineering to the entire life cycle of enterprise data services. The design of the framework will ensure that data pipelines, datasets, reports, dashboards, analytical services, and downstream data products are easily accessible, up-to-date and consistent, visible and consistently operational. The new model contrasts the old model of data-quality which is over-all less concerned with correctness and completeness with the needs of enterprise data capabilities as a production service with reliability objectives that can be measured. The framework is a combination of Service Level Indicators (SLIs), Service Level Objectives (SLOs), error budget, observability, dependency-conscious monitoring, intelligent incident management, automated remediation, and continuous improvement in a single reliability continuum.

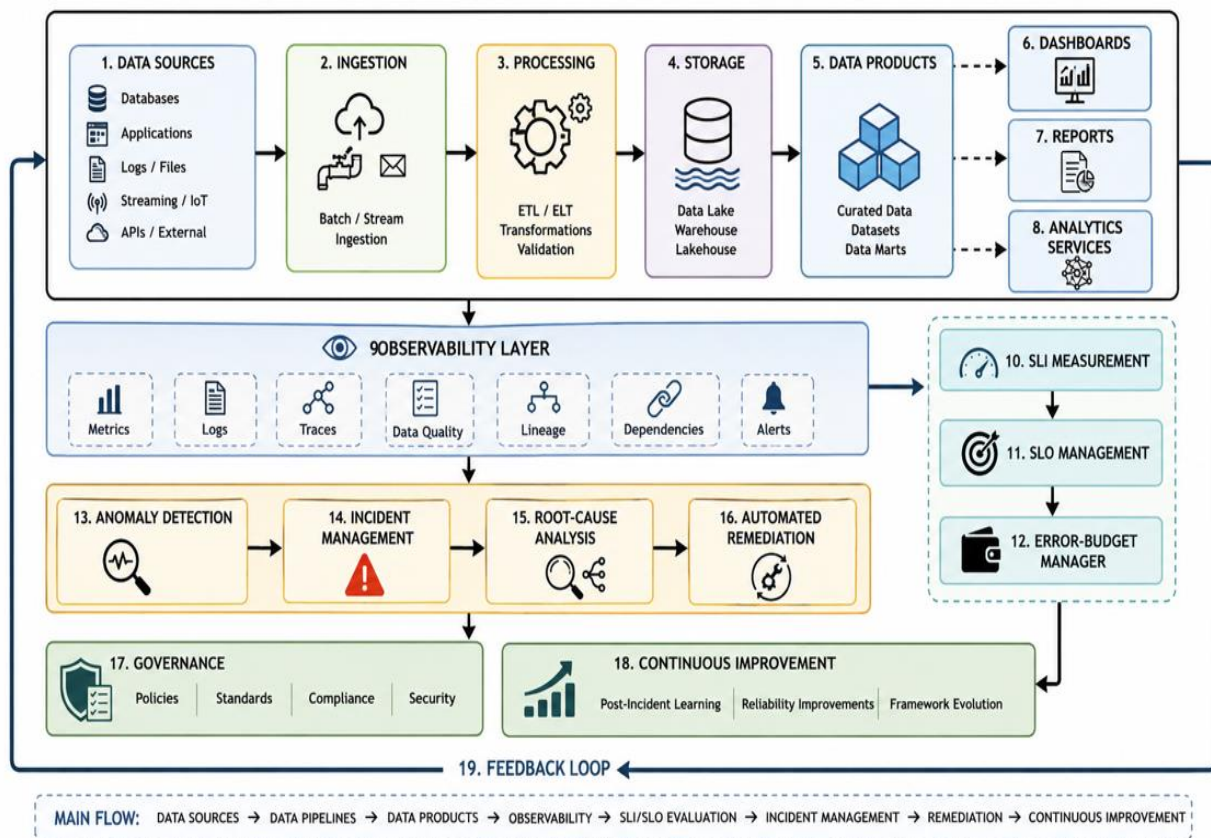


Figure 1- SRE-Driven Reliability Framework for Enterprise Data Platforms

3.1 Framework Architecture

It is an architecture with seven layers, which are interconnected: Data Source and Ingestion Layer, Data Processing Layer, Data Product Layer, Observability and Reliability Measurement Layer, Intelligence and Incident Management Layer, Automated Remediation Layer, and Governance and Continuous Improvement Layer. All these layers are supposed to be used to provide end-to-end reliability-measuring and reliability-improvement mechanism.

Data Source and Ingestion Layer is the initial layer in the Enterprise data ecosystem. It contains transactional databases, enterprise applications, IoT, APIs, cloud applications, external data providers, streaming, file based systems. The data is constantly taken in using either a batch or streaming ingestion pipelines. The source-related measurements, like the success rate of ingestion, availability of sources, latency of ingestion, data arrival, etc. are observed in this layer.

Data Processing Layer: This Data Processing Layer is where the extraction, transformation, validation, enrichment, aggregation and loading operations are found. It can consist of ETL/ELT pipelines, stream processing engines, workflow orchestrators, data warehouses, data lakes, and lakehouses. The reliability within this layer is ascertained by the success when running the pipeline, the processing latency, the frequency of failures, the rate of re-tries and data



completeness. There are also dependency links between pipelines that are maintained to determine the possible downstream effects of failures.

The outputs of the business-facing platform are stored in Data Product Layer, to be made available in datasets, analytics to tables, dashboards and reports, APIs, machine-learning and decision-support software. This layer is of special importance as the availability of infrastructure does not always mean the data-service reliability. The Dashboard can be also available, showing dead or child information. In this way, the framework takes into consideration the technical availability and data specific reliability characteristics.

3.2 Reliability Measurement Using SLIs

The fourth level is the quantitative background of the framework and in this case, it works with the service level indicators (SLIs). SLI is a quantifiable quality of a data service, which represents an operationally stable or consumer-wise reliable operation. There are several types of SLIs defined by the proposed framework.

Pipeline Availability SLI measures the percentage of scheduled pipeline executions completed successfully within the defined operational window:

$$SLI_{availability} = (\text{Total Scheduled Executions} / \text{Successful Executions}) \times 100$$

Data Freshness SLI measures the percentage of data deliveries that satisfy the required freshness threshold:

$$SLI_{freshness} = (\text{Total Expected Deliveries} / \text{Deliveries Within Freshness Target}) \times 100$$

Likewise, SLI Data Quality is a percentage of records or datasets that meet the established or defined quality criteria; and SLI Analytics Availability is the presence of dashboards, reports, APIs, and analytical services. The processing latency, completeness, schema stability and frequency of failed job and the number of incidents can be measured using the additional SLIs.

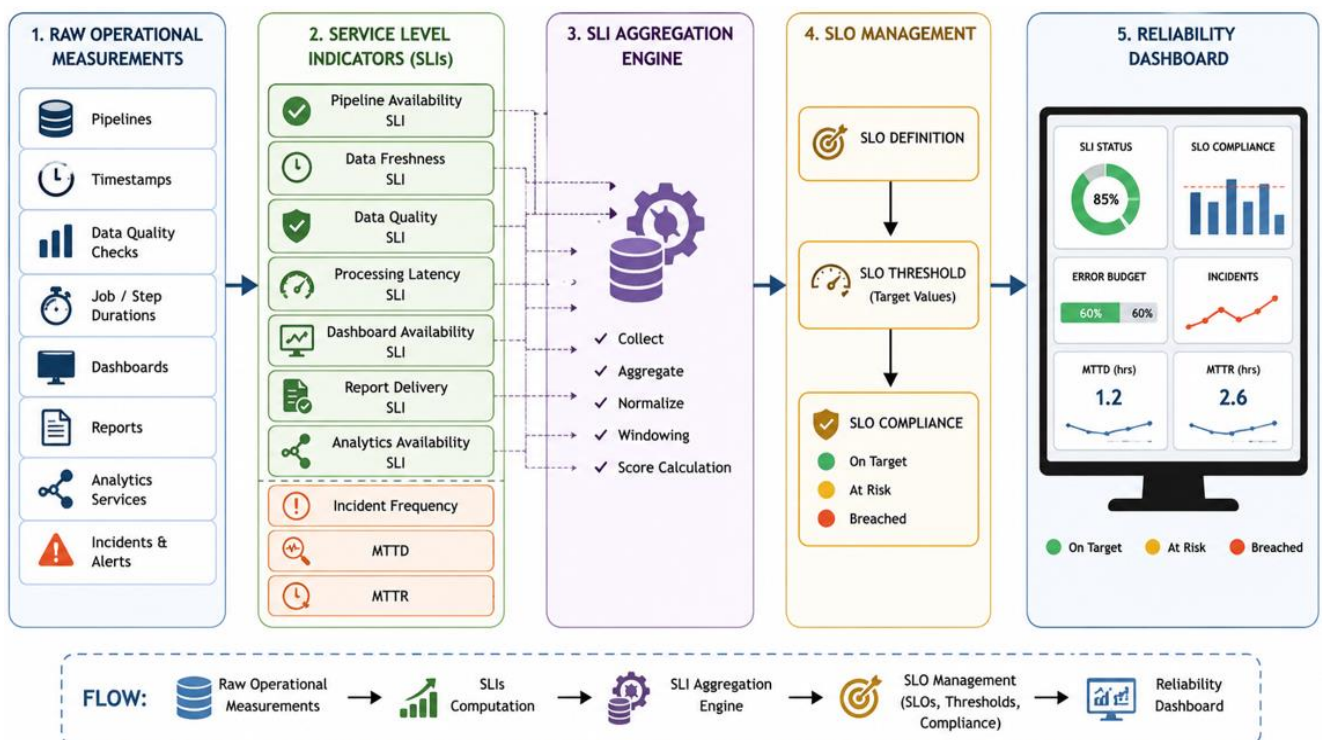


Figure 2- SLI and SLO-Based Reliability Measurement Model

3.3 SLO Management

Services LISs are turned into Service Level Objectives (SLOs) where acceptable reliability objectives of discrete data services are established. These SLOs need to be identified depending on the level of business value rather than implement the same SLOs of all data products. Bearing in mind a concrete case, tougher freshness and availability



requirements can be to a great extent enforced on a real time fraud-detection dataset as compared to a monthly report to the management team.

An SLO could have a 99.9% availability requirement of a crucial service over a pipeline, a 99% need of data-freshness, and a set upper attendance at the processing latency requirement. Every data product includes SLO profile (availability goals, freshness goals, latency goals and quality goals). Violation of SLO will automatically lead to the occurrence of reliability events and might also lead to incident-management procedures.

3.4 Error-Budget Mechanism

The platform also halts and balances platform innovation and production stability with an error budget incorporated in the framework. With an SLO of 99.9, the 0.1 that is remaining is the allowance of reliability in the specified time frame of the measurement. The framework is continuous in estimating the remaining budget and keeping track of how it is used.

With error-budget consumption at low levels, development teams are able to keep on adding new features and pipeline changes. In case the budget is quickly burned, the platform can scale up the monitoring intensity automatically, limit changes that are high risk or focus on reliability engineering efforts. This mechanism poses an objective guideline on whether to focus more on the organization new functionality or stability of operations.

3.5 Observability and Dependency-Aware Monitoring

Continuous processing of logs and metrics, traces, pipeline metadata, lineage data, data-quality outcomes and service-health indicators is achieved using the Observability and Reliability Measurement Layer. Observability will help attain the visibility not only of how a pipeline has failed but also the reason why it has failed and dependencies affected, as well as what business services, which may be impacted, are dependent on it.

There is dependency among data sources, pipelines, datasets, reports, dashboards and analytical applications. The dependency graph is used by the framework to identify the impacted downstream services in case of a failure or delay in an upstream component. This enables the performance of priorities on incidences to be carried out intelligently in accordance to the business impact.

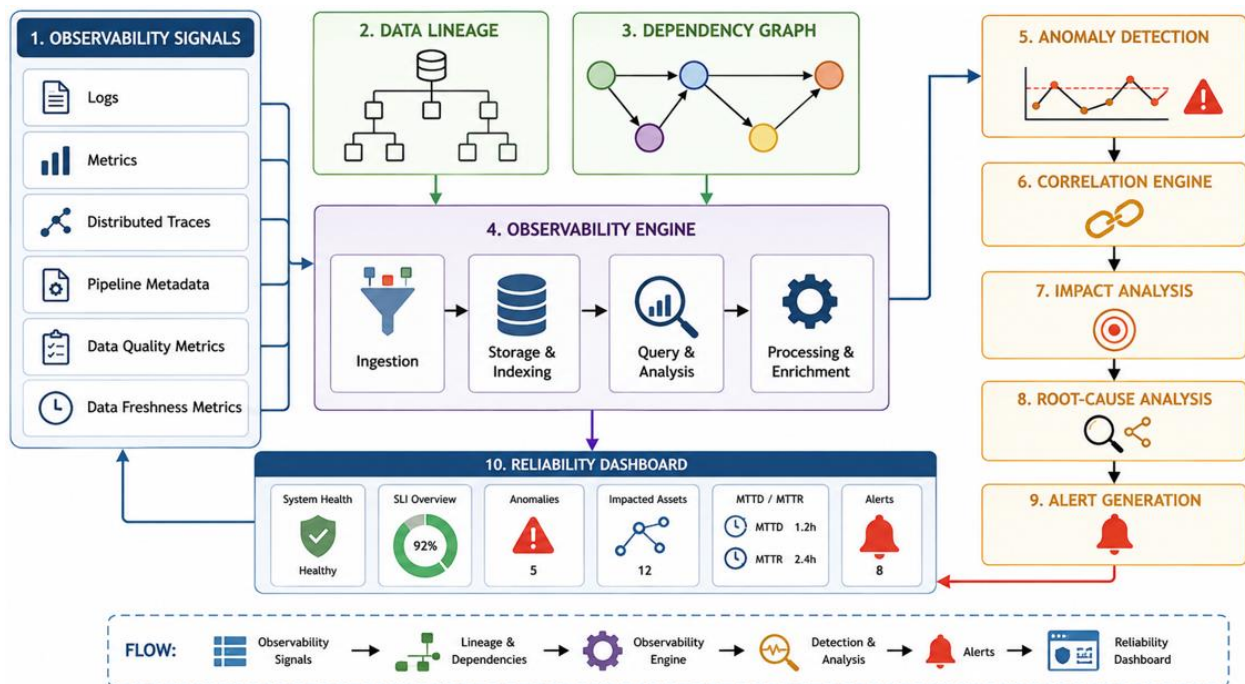


Figure 3- Observability and Dependency-Aware Monitoring Architecture



3.6 Intelligent Incident Management

The reliability events are identified with the help of the Intelligence and Incident Management Layer, and the recurring failure patterns, which are the potential root causes of failure are determined. Anomaly detection based on machine-learning can detect machine-specific strange processing times, changes in data volumes (that were not predicted by the algorithm), violation of freshness, pipeline recurrent failures, or strange resource consumption.

The classification is determined by the severity and services that are affected, SLO that are affected and the error budget remainder of the incident. The top priority and automatic escalation is applied to critical incidents. Miscellaneous timestamps of an incident are registered in the framework to compute Mean Time to Detect (MTTD) and Mean Time to Recover (MTTR). Decrease of these values provides a higher responsiveness of operation.

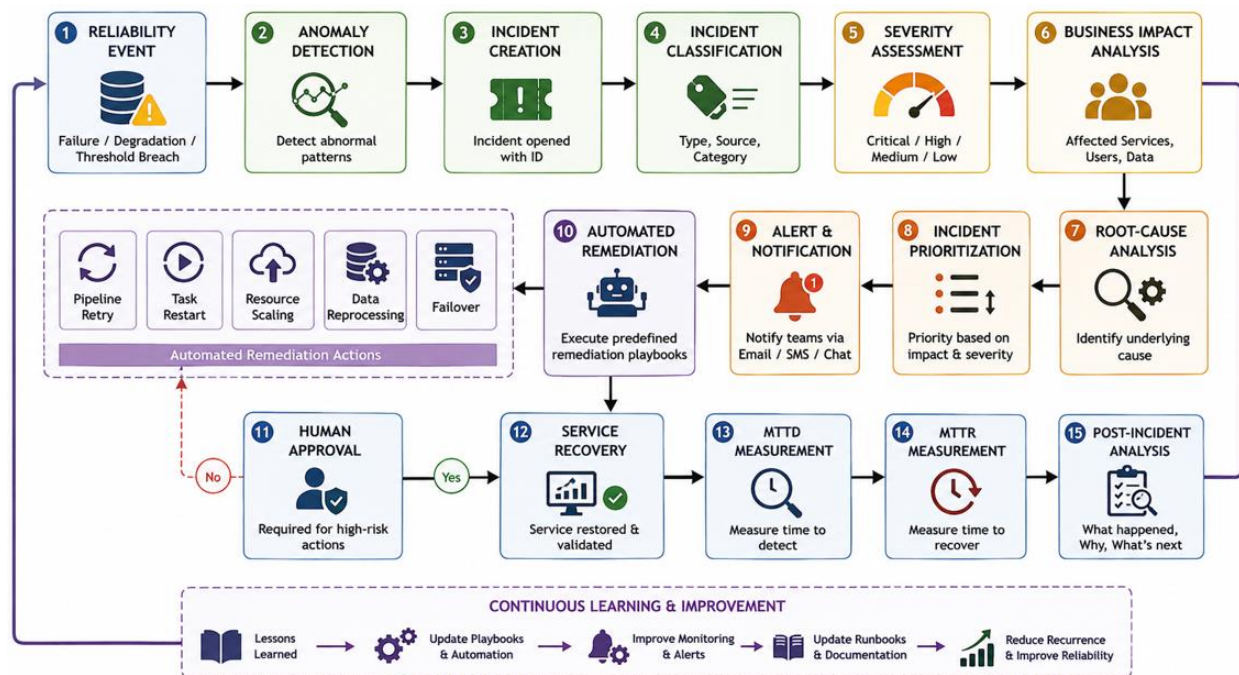


Figure 4- Intelligent Incident Management and Automated Remediation Workflow

3.7 Automated Remediation

Automated Remediation Layer offers the appropriate remedies of the pre-determined failure states. The remediation can either be a restart of the pipeline, re-initiating a task, scaling of resources or re-scheduling of the workflow, re-processing the data, service breakage or alerting depending on the nature of the incident. The automated remediation is supposed to be under set policies to avoid uncontrolled alterations in production systems.

To apply a human-in-the-loop method to incidences of high risk the framework can propose remediation action and an approved engineer endorsing the action. This trade-off on the speed of operational against safety and control.

3.8 Governance and Continuous Improvement

The last tier abandons a Ring of perpetual enhancement. To plot root causes, factors, weaknesses in detection and remediation efficiency, the framework undertakes post-incident analysis, after each major incident. The lessons learned update monitoring policies and SLO targets, policy- and architecture-dependent models of dependencies, policies and pipeline architecture.

The governance systems help in ensuring reliability objectives are adherent to the requirements of the business, security policies, requirements on data quality and regulatory requirement. The markers that can be provided through the reliability reports to the management are: SLO compliance, frequency of occurrence of incident, error-budget consumption, MTTD, MTTR, freshness compliance and pipeline availability.



IV. FRAMEWORK EVALUATION

The proposed SRE-Driven Reliability Framework on Enterprise Data Platforms is experimented with to determine whether it is useful in improving reliability of data-services, operation responsiveness and data availability. The evaluation metrics would consist of the framework capabilities to continuously verify data services to the enterprise to determine breaking of reliability, minimize the time to detect code-recovery and confirm that data is up-to-date and pipeline online. The evaluation methodology also considers the conventional metrics of data-platforms in addition to SRE-focused metrics, such as SLO adherence, error-budget consumption, Mean Time to Detect (MTTD), and Mean Time to Recover (MTTR).

4.1 Evaluation Objectives

The last is to find out whether the proposed structure is able to deliver the additions that can be measured against the traditional data-platform practices which relies on observation. In this measurement, five key targets have been taken into account namely (i) better pipeline availability, (ii) updated data, (iii) incident detectiveness, (iv) incident recovery time and (v) realisation of pre-established goals of SLOs. The other minor objective is to establish whether error-budget monitoring could be used to better prioritize the activities in reliability engineering.

4.2 Experimental Environment

The architecture can be tested based on a representative enterprise data-platform architecture that includes data in batch and streaming data pipelines, data-storage infrastructure, data-transformation infrastructure, analytical data, dashboards, and reporting services. Functional events (data modeling) Model of typical production failures Typical production failures perhaps a model of synthetic or historic operational events may be added: failure of pipeline execution, slows of incoming data, schema evolution, processing bottleneck, dependency failure, and violations of data-quality.

Two scenarios are investigated one standard surveillance (baseline) and the intended SRE-based infrastructures. Common monitoring and alerting of pipelines is part of the baseline, however, the new technique utilizes SLOs, SLIs, error budgets, dependency-aware observability, as well as smart incident detection and automatic remediation.

4.3 Reliability Metrics

The pipelines availability, compliance with the data freshness, compliance with SLO, MTTD and MTTR are the key measurement indicators. Availability of the pipeline is the percentage of the successful pipelines operations that happens with the estimated processing time. Data-freshness compliance is measure of how many of the expected data deliveries meet specified freshness requirements.

SLO compliance is calculated as:

$$\text{SLOCompliance} = (\text{Total Measurement Time} / \text{Time Within SLO}) \times 100$$

MTTD refers to an average of time lag between fulfillment of a reliability event and initiation of the event compared to MTTR which is an average time lag to restoration of service to which the incident under consideration has taken place. Low MTTD and MTTR indicates increasing responsiveness to operation.

4.4 Error-Budget Evaluation

The error-budget consumption is considered to be the estimation of the probability of the framework to operate the effective system of the innovation and reliability balance. Any data service error budget is founded on SLO of a data service. The framework will monitor the budget consumption and inform the consumption where it surpasses the limits which are set beforehand.

To illustrate using a critical pipeline as an example, the 0.1 percentage of the measurement time is the toleration within reliability in the instance of a critical pipeline with an availability SLO of 99.9. This funding funds out rapidly and leads to increased follow up and priorities of reliability improvements. The analysis will discern the possibility of the error-budgeting policies to identify at-risk-services in time before the common jockeys can overpower business users in a scale that would be material.

4.5 Incident Detection and Recovery Evaluation

Control lab test is based on controlled reliability and incidents which test infrastructure of the framework in terms of dealing with incidents. These may be: data-freshness violation, high latency in processing, broken upstream sources, and the broken ingestion jobs and slow transformations. The information of each happenstance is taken as the time of detection, classification accuracy, time of escalation, remedial time and success of recovery.



The proposed framework will enhance the incident response with continuous observability and dependency-based analysis and automatic remediation. A comparative analysis of the proposed framework and the baseline one can prove the decrease of MTTD and MTTR and increase of successful recovery rates.

4.6 Overall Evaluation

The overall evaluation along with the reliability, observability and measures of operation ranks as the overall assessment of the overall performance of the frameworks. Better access to some pipelines and better data-point freshness are indicative of a more reliable data-service, as is the enhanced responsiveness of the operations because of the decrease in MTTD and MTTR. The compliance to SLO and error-budget consumption are some of the other evidences that can be used to determine whether the framework is able to handle reliability in a systematic manner.

It is analysis that provides on the fact that the proposed framework will be able to change passive recovery of failures into proactive, quantified and ever more monitoring SRE-based model of enterprise data activities.

V. CONCLUSION AND FUTURE WORK

5.1 Conclusion

The study introduced a Reliability Framework of Enterprise Data Platforms with SRE to respond to the increased demand of measurable, dynamic and long-term reliability of the contemporary data ecosystems. The framework integrates Site Reliability Engineering and data-platform capabilities of observability, data-quality control, pipeline health, data freshness, dependency management, incident response and automated remediation. The framework makes organizations set quantifiable reliability goals of pipelines, datasets, dashboards, reports and analytical services by defining Service Level Indicators (SLIs) and Service Level Objectives (SLOs). Systematic mechanism offers a balance of both the new development and stability of production through existence of error budgets.

The structure is also dependency sensitive monitored on and intelligent incidence management approaches that reveals the break of reliability and possible impact on a downstream. The MTTD and the MTTR measures give quantitative measurements of incident-detection and recovery. As a result, the suggested strategy will redirect the work of enterprise data towards post-factum work with failures and provide ongoing measuring reliability and its enhancement. The framework offers a single point to enhance the availability of the oil pipes, the freshness of the data, the availability of the analytics, operational resilience and the trust of the enterprise data products.

5.2 Future Work

More studies will be based on the application and testing of the framework on a large-scale production environment, real-world working data, and various types of cloud and hybrid-cloud technology, and solutions. Machine-learning can be used to supplement generative-AI to improve the detection of anomalies, root-cause discoveries, the prioritization of incidents, and fruitful forecasting of failure. The future will also see the creation of even more adaptative SLOs and dynamic error budgets (that vary with workload properties, business criticality profiles and past reliability profiles).

Other significant orientations are the direction in terms of creating a standardized reliability score to allow comparison between a heterogeneous enterprise data product. The impact analysis and the automated remediation can also be further integrated with the integration with data lineage, knowledge graphs and intelligent dependency models. To validate the framework proposed in the future experiments, the framework needs to be compared with the traditional monitoring solutions of the controlled production incidents, and longitudinal measures of SLO compliance, MTTD, MTTR, freshness and availability. Finally, security, privacy, governance and regulatory issues will be incorporated more to achieve a complete framework to reliable, trustworthy and strong enterprise data platforms.

REFERENCES

- [1] M. R. Mekala, "Observability in AI-driven pipelines: A framework for real-time monitoring and debugging," *International Journal of Research in Computer Applications and Information Technology*, vol. 8, no. 1, pp. 686–702, 2025.
- [2] N. Bangad, V. Jayaram, M. S. Krishnappa, A. R. Banarse, D. M. Bidkar, A. Nagpal, and V. Parlapalli, "A theoretical framework for AI-driven data quality monitoring in high-volume data environments," *arXiv preprint arXiv:2410.08576*, 2024.
- [3] S. Patel, R. Gupta, and T. Shaw, "Standardizing AI-driven data quality checks within mesh architectures," *Journal of Data Science and Quality Assurance*, vol. 12, no. 1, pp. 70–85, 2023.



- [4] L. Hummer, S. Green, and M. Lee, “Data lineage techniques for data-intensive applications,” *Data Engineering Review*, vol. 15, no. 4, pp. 210–225, 2023.
- [5] X. Liu and J. Smith, “Real-time AI applications for monitoring data integrity,” *Journal of Data Integrity and Quality*, vol. 5, no. 2, pp. 34–45, 2023.
- [6] International Organization for Standardization, *ISO/IEC 25012:2008—Software Engineering—Software Product Quality Requirements and Evaluation (SQuaRE)—Data Quality Model*. Geneva, Switzerland: ISO, 2008.
- [7] IEEE, *IEEE Standard for an Architectural Framework for the Internet of Things (IoT)*, *IEEE Std 2413-2022*. Piscataway, NJ, USA: IEEE, 2022.
- [8] IEEE, *IEEE Standard for Intercloud Interoperability and Federation (P2302)*. Piscataway, NJ, USA: IEEE, 2021.
- [9] E. Tabassi, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1. Gaithersburg, MD, USA: National Institute of Standards and Technology, Jan. 2023, doi: 10.6028/NIST.AI.100-1.
- [10] IEEE, *IEEE Standard for Transparency of Autonomous Systems*, *IEEE Std P7001-2022*. Piscataway, NJ, USA: IEEE, 2022.
- [11] K. Shah, S. Gami, and A. Trehan, “An intelligent approach to data quality management: AI-powered quality monitoring in analytics,” *International Journal of Advanced Research in Science Communication and Technology*, vol. 4, no. 3, 2024.
- [12] Gartner, *Gartner AI Maturity Model*. Stamford, CT, USA: Gartner, Inc., 2017.
- [13] European Union, *Artificial Intelligence Act*. Brussels, Belgium: European Commission, 2024.
- [14] European Union, *General Data Protection Regulation (GDPR)*. Brussels, Belgium: European Union, 2018.