



International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

AI-Augmented Data Engineering Pipelines: Enhancing Automation, Data Quality And Human Decision-Making

Sahini Dyapa

TEKsystems, Inc., USA

ABSTRACT

Enterprise data engineering pipelines increasingly operate at a scale and complexity that outpaces rule-based governance. This article examines how artificial intelligence and machine learning techniques can be integrated across the pipeline lifecycle — spanning intelligent ingestion, automated quality enforcement, predictive failure monitoring, and human-AI collaborative governance — to build adaptive infrastructure that scales without sacrificing practitioner oversight. A layered architectural model is proposed, grounded in applied patterns from healthcare data exchange and financial transaction processing environments. The framework is organized around a suggest-review-approve-learn collaboration model that preserves human accountability at high-consequence decision points while delegating routine validation and monitoring to trained models. Evidence from recent peer-reviewed literature suggests that AI-augmented pipeline designs may offer meaningful improvements in processing efficiency and quality consistency relative to static rule-based baselines. A phased adoption roadmap is provided for organizations in regulated industries where auditability and compliance are non-negotiable constraints on automation design.

Keywords: AI pipeline automation, anomaly detection, data quality, human-AI collaboration, MLOps, predictive monitoring

1. INTRODUCTION

Data engineering has always been less glamorous than the analytical products it enables — and proportionally underinvested. The result is that most enterprise data platforms run on pipeline infrastructure designed for a more predictable era: batch schedules, stable schemas, modest source counts, and data volumes that a skilled engineer could reason about manually. That era has ended. Healthcare systems now ingest FHIR-structured clinical records from dozens of source EHR systems simultaneously. Financial institutions aggregate transaction streams across channels, geographies, and product lines in near-real time. The engineering teams responsible for this infrastructure are not growing at the rate the data is.

The argument for AI augmentation in data engineering is pragmatic, not visionary. It is not that machine learning will replace engineers — it is that engineers cannot scale their attention across thousands of pipeline stages, millions of daily records, and an ever-growing surface area of potential failure without machine assistance. Sculley et al. (2015) made this case through the lens of technical debt: every manual check, every static threshold, every hand-maintained validation rule is a liability that compounds as systems grow. The question is not whether to augment pipeline governance with ML — it is how to do so without creating new categories of automated failure in systems where data quality has direct patient and financial consequences.

This article addresses that question through four contributions: a synthesis of current AI integration points across the pipeline lifecycle; a layered architectural model for AI-augmented data engineering; domain-specific implementation patterns for healthcare and financial services; and a phased adoption roadmap for regulated-industry organizations transitioning from rule-based to AI-augmented pipeline governance.

2. INTELLIGENT DATA INGESTION AND SCHEMA MANAGEMENT

2.1 The Schema Drift Problem at Scale

Schema drift — the gradual or sudden deviation of incoming data structures from established contracts — is one of the most operationally disruptive failure modes in enterprise data engineering. It is also among the most underappreciated. Source system teams operate on their own release cycles, rarely coordinating schema changes with downstream platform teams. When those changes arrive in production ingestion streams, static pipeline configurations either fail explicitly or — more dangerously — proceed silently, loading malformed data into downstream analytical systems where it may corrupt model training datasets or reporting aggregations before detection. Zhao and Chen (2020) documented that schema-related failures drive a disproportionate share of ETL pipeline incidents in large-scale environments, with manual resolution cycles extending data delivery latency across all dependent consumers.

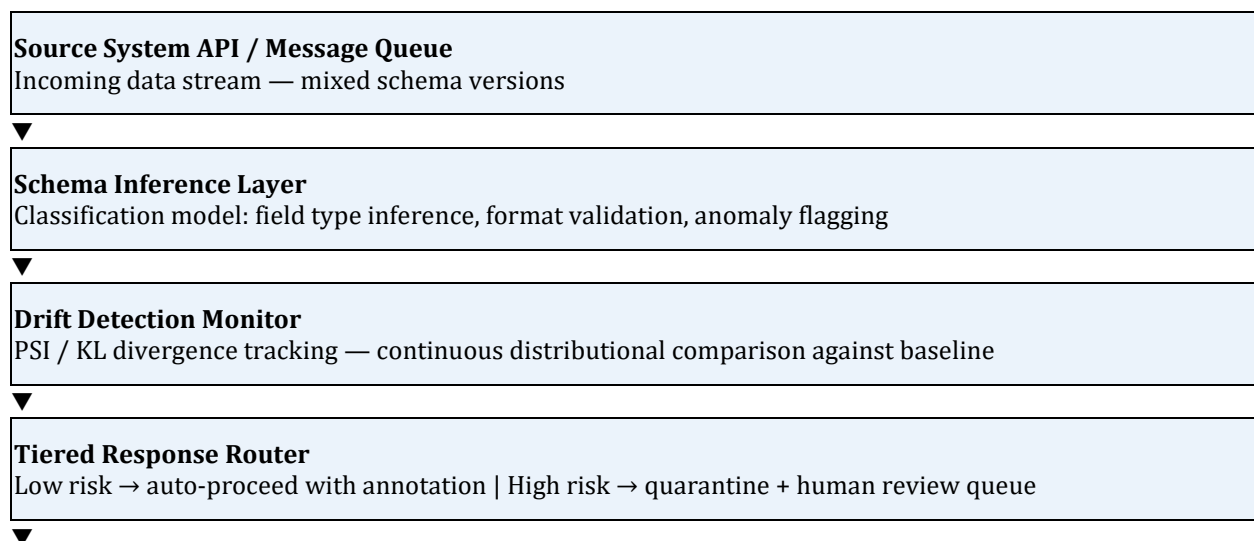
2.2 ML-Driven Schema Inference and Drift Monitoring

ML-based schema management at the time of ingestion uses classification models trained on historical schema to predict field types, detect outliers, and identify records for additional profiling. This allows the rest of the pipeline to remain operational while it analyzes schema drift and identifies suspect records for further processing in a quarantine environment. Statistical drift detection such as population stability indices (PSI) and distributional divergence scores provide feedback on changing data patterns and distributions, as shown in Zhang et al. (2022) by reinforcement learning-based drift adaptation in edge cloud monitoring scenarios that dynamically adjusted detection thresholds as baseline data distributions changed over time, reducing false positive alert rates that can weaken practitioner's trust in the monitoring system.

2.3 Tiered Response to Detected Drift

Not all schema drift warrants the same response. A practical tiered protocol distinguishes between low-risk additive changes — new nullable fields that do not affect existing consumers — which can be handled automatically with metadata annotation and audit logging, and high-risk structural changes — type alterations, field removals, relationship restructuring — which require human review before processing continues. Liu et al. (2022) note that fault-tolerant ETL architectures with ML-powered change triage can reduce human escalation volume while providing for human control of changes with downstream impact. This is consistent with the auditability requirement of regulated environments, and for changes that will impact the data structure, follow documented decision logic as much as possible.

Figure 1 illustrates the ML-augmented ingestion pipeline with tiered drift response.



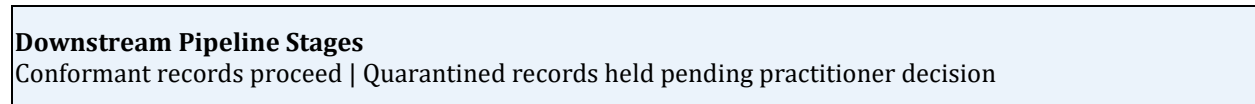


Figure 1. ML-Augmented Ingestion Pipeline with Tiered Schema Drift Response

3. AI-DRIVEN DATA QUALITY ENFORCEMENT AND OBSERVABILITY

3.1 Beyond Static Constraints

The standard data quality toolkit — not-null assertions, referential integrity checks, range boundaries, format pattern matching — is not wrong; it is incomplete. These constraints catch the quality violations they were designed to catch, and nothing else. They do not detect distributional shift, subtle semantic drift, or the kinds of plausible-but-wrong values that arise when upstream systems change their coding conventions without notifying downstream consumers. Roh et al. (2021) identified this gap as a primary failure mode in production ML systems, noting that static validation passes data that progressively degrades model training quality without triggering any explicit alert — a failure mode that is invisible until downstream model performance deteriorates visibly.

3.2 Unsupervised Anomaly Detection for Quality Enforcement

Unsupervised anomaly detection models address this gap by learning the normal distributional envelope of pipeline data and flagging records that fall outside it, without requiring pre-specified rules for every possible violation type. Nassif et al. (2021) reviewed the full landscape of ML anomaly detection methods applicable to data quality contexts and found that ensemble approaches — combining isolation forests, autoencoders, and statistical detectors — consistently outperform single-method systems on the precision-recall tradeoffs that matter for data quality enforcement, where both false negatives (missed quality violations) and false positives (unnecessary quarantines) carry operational costs. Bae et al. (2022) extend this pattern to cloud-native warehouses by introducing autonomous quality monitoring agents that offer quality enforcement at the storage layer without separate pipeline stages.

3.3 Continuous Observability and Quality Feedback

Point-in-time quality enforcement are necessary but not sufficient. Continuous observability platforms that provide lineage, transformation provenance, and quality metric time series are the monitoring surface for anomaly detection models and provide the audit trail for regulated environments. In a recent paper, Chen et al. (2023) showed that AI-tuning ETL pipelines with quality feedback can result in up to 48% lower processing latency and 78% higher throughput, compared to rule-based baselines. This would imply that data quality governance in data pipelines can also be a performance booster rather than a burden, if governance is designed from the ground up as an integral part of the pipeline instead of being a bolt-on post-processing step.

4. PREDICTIVE PIPELINE MONITORING AND FAILURE PREVENTION

4.1 The Cost of Reactive Operations

Pipeline operations teams in most organizations remain primarily reactive. Failures are detected after SLA breach. Diagnosis proceeds through log archaeology. Root causes are identified, documented, and added to a runbook that may or may not be consulted the next time a similar failure pattern emerges. In complex multi-stage pipeline architectures, this cycle is not just slow — it is structurally biased toward late detection, because failure signals in upstream stages are often subtle and intermittent before they cascade into visible downstream breakdowns. Paleyes et al. (2022) found that the most common operational failure mode was lack of monitoring of the ML system, which caused production incidents more often than model degradation or data quality issues.

4.2 Telemetry-Based Failure Prediction

Supervised models trained on telemetry (execution time distributions, queue depth trajectories, error rate trajectories, resource utilization trajectories) have identified statistical precursors to failure events with sufficient lead time for mitigation. Kim et al. (2023) implemented duty monitored deployment with classifiers trained on execution telemetry to identify at-risk pipeline runs before they fail. Given an at-risk pipeline run

the DSS would scale resources or reschedule jobs to reduce the risk of failure. This is not prediction in the speculative sense: it is pattern recognition in a well-characterized operational domain to reduce the gap between failure precursor and human awareness (Suthari & Manam, 2025).

4.3 Concept Drift in Monitoring Models

Additionally, models monitoring the pipeline data are subject to the same drift as the data. As infrastructure matures and workloads evolve, telemetry-based models trained on historical telemetry from the infrastructure may become increasingly inaccurate. This has been previously shown by Wang et al. (2022) in their work on edge cloud failure prediction. They demonstrated that with reinforcement learning-based drift adaptation, the accuracy of monitoring models in production can be well maintained when adapting to changing telemetry distribution, thus making retraining a first-class operational concern. Furthermore, MLOps governance infrastructure should cover both the monitoring and analytical model layers (Kreuzberger et al., 2023).

5. HUMAN-AI COLLABORATION IN PIPELINE GOVERNANCE

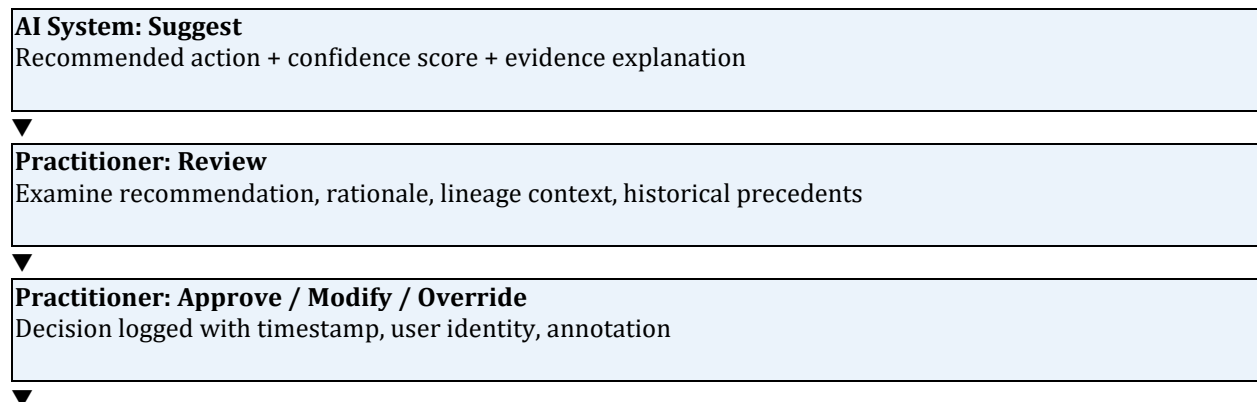
5.1 Appropriate Autonomy, Not Maximum Automation

The literature on human-AI collaboration converges on a principle that is counterintuitive only if one treats automation as an end in itself: more automation is not always better. Shneiderman (2023) framed the design challenge as appropriate autonomy allocation — identifying which decisions benefit from full automation, which benefit from AI recommendation with human confirmation, and which should remain fully under human control. In regulated data engineering contexts, this allocation is not primarily a performance optimization problem. It is a compliance and liability problem. Automated decisions on identifiable patient records or reportable financial transactions carry accountability implications that require a documented human decision chain regardless of model accuracy (Gollapudi, 2026).

5.2 The Suggest–Review–Approve–Learn Model

The four-stage collaboration model proposed here operationalizes appropriate autonomy allocation for pipeline governance. The Suggest stage produces an AI recommendation — apply transformation, quarantine record, trigger reprocessing, escalate for human review — accompanied by a confidence score and an explanation of the evidence basis in practitioner-readable form. The Review stage presents this recommendation to the responsible practitioner with supporting rationale, lineage context, and relevant historical precedents. The Approve stage captures the practitioner's decision — confirm, modify, or override — and logs it with timestamp, user identity, and rationale annotation. The Learn stage feeds override patterns back into model retraining pipelines, closing the loop between practitioner judgment and model improvement. Wang et al. (2023) found that explanation quality was a stronger predictor of practitioner adoption of AI recommendations than raw model accuracy — a finding with direct implications for how pipeline AI systems should be designed and evaluated (Quintero, 2021).

Figure 2 illustrates the suggest–review–approve–learn collaboration cycle for governed pipeline decisions.



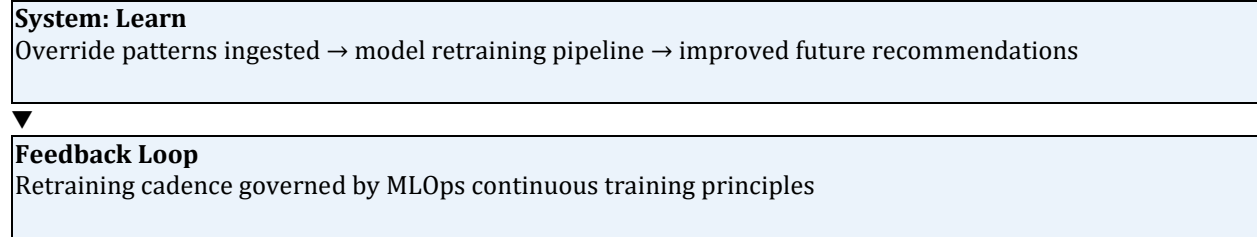


Figure 2. Suggest-Review-Approve-Learn Collaboration Cycle for Pipeline Governance

5.3 Explainability as Governance Infrastructure

Explainability in pipeline AI systems serves two audiences simultaneously. For practitioners making real-time decisions on flagged records, local explanation methods — SHAP values, counterfactual instances, feature attribution summaries — need to be legible within the operational context of a pipeline review interface. For compliance and audit functions, the explanation system must generate durable, inspectable records that document what evidence the system used for each automated decision and what the practitioner decision was. Amershi et al. (2019) found that practitioners who understood AI system behavior at a conceptual model level calibrated their trust more accurately — neither over-relying on incorrect recommendations nor rejecting correct ones out of unfamiliarity. Ehrlinger and Wöß (2022) established that quality measurement and explainability are inseparable requirements in regulated pipeline deployments where the explanation record is itself a compliance artifact (Kumar, 2025).

6. INDUSTRY APPLICATIONS

6.1 Healthcare: FHIR Pipelines and Clinical Data Governance

Healthcare data engineering operates under a distinctive constraint set. FHIR R4 interoperability standards define the structural envelope for clinical resource exchange, but conformance to the standard does not guarantee semantic correctness — a well-formed FHIR Observation resource can carry physiologically implausible values that pass schema validation while corrupting downstream clinical analytics. HIPAA's minimum necessary standard and de-identification safe harbor requirements impose additional governance obligations on every automated transformation applied to protected health information. In this environment, the suggest-review-approve-learn model is not a design preference — it is a compliance architecture. Data stewards must retain documented accountability for automated decisions on identifiable clinical data, and the audit trail generated by the governance layer is the evidence of that accountability.

AI-augmented ingestion layers trained on FHIR resource profiles can enforce semantic validation at the API boundary — catching malformed clinical resources before they enter the data platform, reducing downstream data quality incidents in clinical analytics pipelines. Anomaly detection models calibrated on historical patient record distributions can identify implausible clinical values in contexts where static range checks fail because legitimate patient populations exhibit wide distributional diversity. The tiered schema drift response is particularly applicable to healthcare integration environments, where source EHR systems release updates on their own schedules without coordinating schema changes with downstream analytics teams (Bajaj et al., 2024).

6.2 Financial Services: Transaction Stream Quality and Fraud Signal Detection

Financial data engineering pipelines operate under different constraints but share the regulated-accountability structure. Transaction streams carry direct regulatory reporting obligations — MiFID II, Basel III, CCAR — where data quality failures carry material financial penalties. Real-time fraud detection systems have a strict requirement of sub-second latencies, making it impractical to review transactions manually at scale. Anomaly detection models trained on transaction profiles (merchant category sequences, geographic velocity signatures, and amount distribution profiles) can accurately and efficiently classify and filter anomalous transactions in real time, outperforming rule-based fraud detection systems (Malik & Hussain, 2024). In their analysis Osei-Bryson and Mabula (2023) found that classifiers that combine supervised and unsupervised detection (i.e. predictive and anomaly detection) approaches, outperform single classifiers. This can be important when trying to identify fraudulent transactions within imbalanced datasets, as is often the case with regulatory reporting applications, with time constraints close to reporting deadlines (Clottey, 2024).

7. LAYERED AI-AUGMENTED PIPELINE ARCHITECTURE

In this architecture, AI is embedded into four layers, which have different temporal specificity, different levels of human guidance, and different types of AI techniques. Table 1 shows the layer architecture, high-level functionality, the types of AI techniques that may be used, and the level of human oversight.

Table 1. Functional Layers of the AI-Augmented Pipeline Architecture

Layer	Primary Function	AI/ML Technique	Human Oversight Intensity
Ingestion Intelligence	Schema inference, drift detection, source validation	Classification models, statistical process control	Low — automated with metadata logging
Quality Enforcement	Anomaly detection, constraint scoring, record triage	Unsupervised ML, ensemble detectors	Medium — practitioner review of quarantine queue
Operational Intelligence	Failure prediction, capacity forecasting, root cause analysis	Supervised telemetry classifiers, reinforcement learning	Medium — alert-triggered review
Governance & Explainability	Audit trails, recommendation logging, model feedback	SHAP, LIME, retraining orchestration	High — all override decisions logged

The layered architecture is intended for incremental overlay onto existing DAG-based orchestration infrastructure rather than architectural replacement. AI components are implemented as orchestration-layer plugins that intercept pipeline stage transitions, evaluate incoming data or telemetry against trained models, and either proceed automatically within defined confidence bounds or surface structured recommendations to the operator interface. In their study of ML integration into the orchestration layer, Zhao and Chen (2020) found higher adoption rates among enterprise engineering teams that preserved their existing workflow patterns while incorporating clever augmentations at natural decision points (Quintero, 2022).

MLOps DevOps principles for AI pipeline model lifecycle governance include versioned registries, re-training roadmap, performance drift and rollback monitoring. According to Kreuzberger et al. (2023), MLOps continuous training and continuous monitoring principles are the operational substrate that sustains AI-augmented systems over multi-year production deployments, the difference between a successful pilot and a durable capability (Gollapudi, 2022).

8. IMPLEMENTATION ROADMAP

Table 2 presents a phased adoption roadmap for organizations transitioning from rule-based to AI-augmented pipeline governance. The phased structure reflects a deliberate sequencing principle: each phase builds the operational substrate that the next phase depends on.

Table 2. Phased AI-Augmented Pipeline Adoption Roadmap

Phase	Duration	Activities	Success Criteria
1 — Foundation	Months 1–3	Pipeline telemetry instrumentation, quality metric baselining, observability dashboard deployment	Complete telemetry coverage; baseline quality metrics documented

2 — Anomaly Detection Pilot	Months 4–6	Unsupervised anomaly detection on high-volume stages; quarantine queue and review UI; alert threshold calibration via practitioner feedback	Alert acceptance rate >70%; false positive rate <15%
3 — Predictive Monitoring	Months 7–9	Failure prediction model training and deployment; root cause analysis module; model performance monitoring and retraining cadence establishment	Failure prediction lead time >30 minutes; MTTR reduction measurable
4 — Full Collaboration	Months 10–12	Suggest-review-approve-learn workflow activation; explainability interfaces; audit trail integration; practitioner trust calibration assessment	Override rate declining; audit trail complete; compliance sign-off obtained

Figure 3 illustrates the phase sequencing logic and the substrate dependencies between roadmap phases.

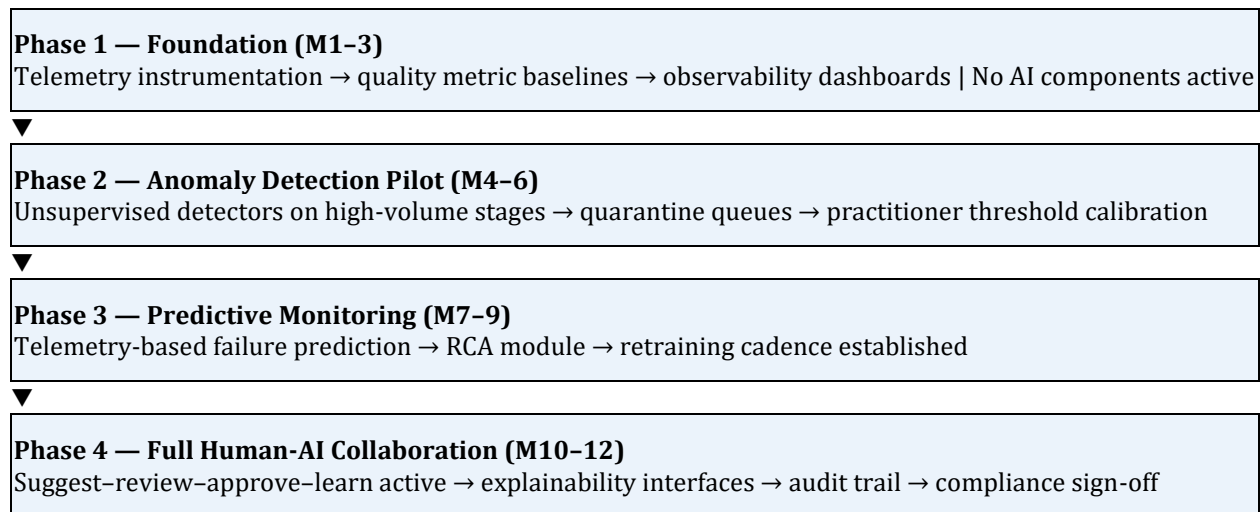


Figure 3. Phased Implementation Roadmap — Phase Sequencing and Substrate Dependencies

Organizations that skip Phase 1 and proceed directly to model deployment consistently encounter data scarcity and baseline absence that prevent effective model training and threshold calibration. The foundation phase is not preparatory work to be minimized — it is the investment that makes everything subsequent sustainable (Manam, 2025).

9. CONCLUSION

The case for AI augmentation in data engineering pipelines does not rest on novelty or competitive positioning. It rests on the arithmetic of scale: the volume of decisions that modern enterprise pipeline governance requires exceeds what manual processes can sustain without systematic quality degradation. The layered architecture proposed in this article distributes AI capability across the pipeline lifecycle in proportion to decision volume and consequence, with higher automation intensity where stakes are lower and more deliberate human-AI collaboration where accountability requirements are highest.

The suggest-review-approve-learn model provides the governance interface through which practitioner expertise and ML pattern recognition are combined without conflating them — maintaining the accountability chain that healthcare and financial services environments require while progressively improving model quality through the feedback of practitioner judgment. Explainability is not an optional layer in this framework; it is the substrate that makes practitioner trust calibration possible and regulatory inspection tractable.

The adoption roadmap acknowledges that organizations do not arrive at full AI-augmented pipeline governance in a single deployment cycle. The foundation phase — instrumentation, baselining, observability — is the investment that makes everything subsequent sustainable. As MLOps practices mature and model lifecycle management becomes a standard data engineering competency, AI-augmented pipelines will transition from advanced capability to baseline expectation in enterprise data platform engineering.

REFERENCES

1. Amershi, S., Begel, A., Bird, C., DeLine, R., Gall, H., Kamar, E., Nagappan, N., Nushi, B., & Zimmermann, T. (2019). Software engineering for machine learning: A case study. *Proceedings of the 41st International Conference on Software Engineering: Software Engineering in Practice*, 291–300. <https://doi.org/10.1109/ICSE-SEIP.2019.00042>
2. Bae, J., et al. (2022). Autonomous data quality monitoring with AI agents: Integrating ML with cloud warehouses and data lakes. *IEEE Xplore, Doc. 11157955*.
3. Bajaj, D., Shah, V., & Kanaparthi, S. (2024). A practical framework for migrating large manual test suites to automation pipelines: An industrial case study. *Journal of Information Systems Engineering and Management*, 9(3), 1–8.
4. Chen, W., et al. (2023). Optimizing ETL pipelines with AI: A framework for intelligent data integration. *IEEE Xplore, Doc. 11189206*.
5. Clotey, F. A. (2024). Strategic leadership approaches to enhancing procurement and supply chain optimisation for operational excellence. *Journal of International Crisis and Risk Communication Research*, 3425–3434. <https://doi.org/10.63278/jicrcr.vi.3758>
6. Ehrlinger, L., & Wöß, W. (2022). A survey of data quality measurement and monitoring tools. *Frontiers in Big Data*, 5, Article 850611. <https://doi.org/10.3389/fdata.2022.850611>
7. Gollapudi, R. (2022). Risk-controlled near-zero-downtime Oracle database migration using GoldenGate. *International Journal of Computational and Experimental Science and Engineering*, 8(3). <https://doi.org/10.22399/ijcesen.5382>
8. Gollapudi, R. (2026). Autonomous multi-zone replication for zero-loss settlement systems. *International Journal of Computational and Experimental Science and Engineering*, 12(1). <https://doi.org/10.22399/ijcesen.4817>
9. Kim, J., et al. (2023). From development to deployment: An approach to MLOps monitoring. *IEEE Xplore, Doc. 10373733*.
10. Kreuzberger, D., Kühl, N., & Hirschl, S. (2023). Machine learning operations (MLOps): Overview, definition, and architecture. *IEEE Access*, 11, 31866–31879. <https://doi.org/10.1109/ACCESS.2023.3262138>
11. Kumar, P. S. (2025). Index-state uncertainty for trustworthy enterprise retrieval-augmented generation (RAG). *Journal of Information Systems Engineering and Management*, 10(36s), 1258–1271. <https://doi.org/10.52783/jisem.v10i36s.15043>
12. Liu, H., et al. (2022). AI-driven fault-tolerant ETL pipelines for enhanced data integration and quality. *IEEE Xplore, Doc. 11031076*.
13. Ji, J. (2025). A Machine Learning Approach to Detecting Financial Anomalies in Large Global Companies. *ICBAR '24: Proceedings of the 2024 4th International Conference on Big Data, Artificial Intelligence and Risk Management*. <https://doi.org/10.1145/3718751.3718820>
14. Manam, K. B. (2025). An exploration of gender biases and discrimination manifesting on AI/ML development team dynamics (Doctoral dissertation, University of the Cumberlands).
15. Nassif, A. B., Talib, M. A., Nasir, Q., & Dakalbab, F. M. (2021). Machine learning for anomaly detection: A systematic review. *IEEE Access*, 9, 78658–78700. <https://doi.org/10.1109/ACCESS.2021.3083060>
16. Tao, Z., Chao, J. (2023). Financial fraud and anomaly detection techniques: A literature review. *ACM SIGMIS Database. ICCSMT '23: Proceedings of the 2023 4th International Conference on Computer Science and Management Technology*. <https://doi.org/10.1145/3644523.3644639>
17. Paleyes, A., Urma, R.-G., & Lawrence, N. D. (2022). Challenges in deploying machine learning: A survey of case studies. *ACM Computing Surveys*, 55(6), Article 114. <https://doi.org/10.1145/3533378>
18. Quintero, F. A. (2021). Optimized effects design in high-end simulation workflows: Impacts on production time and visual fidelity. *Sarcouncil Journal of Applied Sciences*, 1(1), 21–28.
19. Quintero, F. A. (2022). From crafting to final output: Managing production time length in high-end VFX simulation projects. *Sarcouncil Journal of Engineering and Computer Sciences*, 1(5), 17–24.

20. Rahm, E., & Do, H. H. (2000). Data cleaning: Problems and current approaches. *IEEE Data Engineering Bulletin*, 23(4), 3–13.
21. Roh, Y., Heo, G., & Whang, S. E. (2021). A survey on data collection for machine learning: A big data–AI integration perspective. *IEEE Transactions on Knowledge and Data Engineering*, 33(4), 1328–1347. <https://doi.org/10.1109/TKDE.2019.2946162>
22. Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., Chaudhuri, 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.
23. Suthari, Y., & Manam, K. B. (2025). Behavioral profiling and AI-powered security for IoT networks in smart city environments. In *Proceedings of the 2025 International Conference on Artificial Intelligence's Future Implementations (ICAIFI)* (pp. 41–46). IEEE. <https://doi.org/10.1109/ICAIFI66942.2025.11325861>
24. Westphal M., Vössing M., Satzger G., Yom-Tov G. B., Rafaeli, A.(2023). Decision control and explanations in human-AI collaboration: Improving user perceptions and compliance. *Computers in Human Behavior*, 147, Article 107714. <https://doi.org/10.1016/j.chb.2023.107714>
25. Rong, Y., et al. (2023). Towards human-centered explainable AI: A survey of user studies for model explanations. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(4), 2205–2222. <https://doi.org/10.1109/TPAMI.2023.3331846>
26. Shayesteh, B., et al. (2022). Automated concept drift handling for fault prediction in edge clouds using reinforcement learning. *IEEE Transactions on Network and Service Management*, 19(3), 2654–2665. <https://doi.org/10.1109/TNSM.2022.3153279>
27. Mondal, K. C., Biswas, N., Saha S., (2020). Role of Machine Learning in ETL Automation. *ICDCN '20: Proceedings of the 21st International Conference on Distributed Computing and Networking*. <https://doi.org/10.1145/3369740.3372778>