

Original Article

Data-Centric Refactoring: Techniques for Improving Model Quality via Codebase Changes

Dr. Fatou Diop

Department of Sociology, Dakar Social Sciences University, Senegal.

Received: 03-02-2019

Revised: 03-27-2019

Accepted: 04-02-2019

Published: 04-04-2019

ABSTRACT

Machine learning systems often fail to reach optimal performance not because of inadequate model architectures, but due to poorly structured data processing pipelines hidden within the codebase. Data-centric refactoring aims to improve model quality through systematic restructuring of code elements responsible for data collection, preprocessing, transformation, validation, and feature engineering. This paper introduces a comprehensive taxonomy of data-centric refactoring strategies, investigates their application across ML-driven software projects, and evaluates their impact on model accuracy, robustness, maintainability, and reproducibility. By bridging software refactoring principles with data-centric AI practices, the proposed framework demonstrates that code-level improvements to data handling routines can yield substantial gains in model performance while reducing technical debt. Experimental results show that systematically refactoring data pipelines leads to more reliable features, reduced noise propagation, and improved generalization. The findings position data-centric refactoring as a key discipline for modern ML engineering, enabling scalable, interpretable, and production-ready models.

KEYWORDS

Data-Centric AI, Software Refactoring, ML Pipelines, Feature Engineering, Technical Debt, Data Quality, Codebase Optimization, Mlops, Model Reliability, Data Preprocessing.

1. INTRODUCTION

1.1. Background and motivation

Machine learning (ML) has shifted from a purely research-driven activity to an engineering discipline embedded in production software systems; models are now components in larger codebases that handle ingestion, cleaning, transformation, feature construction, training, validation, and deployment. This integration exposes ML to the same software-quality problems that traditional applications have always faced—fragile, duplicated, or poorly documented code, hidden assumptions, and ad-hoc fixes that accumulate over time. Unlike conventional software, however, faults in data-handling code propagate into model behavior in subtle, often non-deterministic ways: a tiny change in preprocessing logic or an unnoticed mismatch in feature scaling can degrade model generalization, introduce label leakage, or produce reproducibility nightmares. The motivation for this paper is to treat those data-related code issues as first-class engineering problems: to show that deliberate, disciplined refactoring of the codebase that supports data pipelines can measurably improve model quality, reduce costly debugging cycles, and lower long-term maintenance costs. By articulating techniques, taxonomies, and evaluation practices, we aim to bridge the gap between software refactoring traditions and modern data-centric ML engineering.

1.2. Gap in current ML engineering practices

Contemporary ML engineering often emphasizes model architecture, hyperparameter tuning, and compute—what is commonly called “model-centric” optimization—while under-investing in the code that prepares and manages the data the model consumes. This imbalance creates a reality in which significant model errors, drift, and brittle production behavior stem not from the chosen algorithm, but from inconsistent preprocessing steps, duplicated transformations across scripts, untracked feature derivations, and fragile dataset versioning. Moreover, standard software-engineering refactoring techniques rarely address properties unique to ML systems, such as feature lineage, statistical assumptions, and data distribution shifts. Tooling and best practices for detecting and resolving “data code smells” are immature; teams often lack metrics that relate codebase changes to downstream model effects, so refactoring work is undervalued or postponed. The gap, therefore, is two-fold: a conceptual gap in mapping refactoring ideas to data-centric problems, and a practical gap in tools, workflows, and evaluation metrics that demonstrate the payoffs of code-level changes on model performance and reliability.

1.3. Importance of data-centric approaches over model-centric optimization

Data-centric approaches reorient priorities from tweaking increasingly complex models toward improving the quantity, quality, and consistency of the data fed into those models. In many industrial and research applications, incremental gains from larger or deeper models are dwarfed by improvements achievable through better feature engineering, consistent preprocessing, and cleaner labels. Data-centric refactoring is complementary to, and sometimes more cost-effective than, model-centric strategies: it reduces noise and bias at the source, makes feature behavior predictable across environments, and prevents accidental leakage that inflates offline metrics but collapses in production.

Additionally, focusing on data quality through codebase changes improves transparency and traceability—attributes critical for debugging, regulatory compliance, and ethical AI—because it anchors dataset provenance and transformations in well-structured, testable code. The crux is that if the data pipeline is reliable, smaller and simpler models often achieve comparable or superior real-world performance with improved interpretability and maintainability.

1.4. Research objectives and contributions

This paper sets out with four interrelated objectives: (1) to define and formalize the concept of data-centric refactoring and map it to recognizable software-engineering patterns; (2) to provide a taxonomy of data-related code smells and concrete refactoring techniques that address them; (3) to propose an implementation framework and measurable evaluation criteria that link codebase changes to model quality and maintainability; and (4) to empirically demonstrate the approach via case studies or experiments that quantify improvements in model performance, robustness, and developer productivity. The contributions therefore include a conceptual framework that unites refactoring and data-centric AI, a catalog of techniques and triggers for when to apply them, practical guidance for integrating refactoring into MLOps pipelines, and empirical evidence that validates the efficacy of these interventions. Together, these artifacts aim to equip practitioners with actionable methods and researchers with a structured agenda for advancing tooling and studies in this space.

1.5. Paper organization

To guide the reader through this material we begin with a review of related work that situates our proposal within software refactoring, data-centric AI, and technical-debt research. Next, we develop the conceptual foundations of data-centric refactoring—defining scope, identifying key properties of dataflows in ML systems, and introducing a taxonomy of data-code smells. The following section enumerates specific refactoring techniques categorized by pipeline stage (ingestion, preprocessing, feature engineering, validation, training) and describes tool and automation support. We then present a proposed framework for implementing refactoring at scale, including workflow steps, integration points with CI/CD, and evaluation metrics. After that we report experimental evaluations and case studies that compare baseline and refactored systems, analyze results, and discuss threats to validity. We close with a discussion of implications, limitations, and directions for future work. This structure is designed to move from theory and motivation to practical methods and evidence, enabling both conceptual understanding and immediate application.

2. RELATED WORK

2.1. Traditional software refactoring techniques

Software refactoring is a mature field concerned with modifying code structure without changing external behavior to improve attributes such as readability, modularity, extensibility, and testability. Classic refactorings—extract method, rename variable, introduce parameter object, remove duplication—are documented in catalogs and tool-supported in many IDEs; empirical studies have shown that regular refactoring reduces defect rates and eases maintenance. However, those techniques

were developed with deterministic control-flow and stateful business logic in mind, where correctness is judged primarily by functional tests. When applied to data pipelines and ML code, the behavioral invariants differ: transformations that preserve API contracts may nevertheless alter the statistical properties of the data in ways that affect model training. Thus, traditional refactoring provides valuable discipline and techniques, but needs adaptation to capture non-functional, statistical invariants that are crucial for ML systems – an area in which literature is emerging but still sparse.

2.2. Data-centric AI paradigm

The data-centric AI paradigm advocates prioritizing improvements in data – quality, labeling, representativeness, and pipeline integrity rather than relentlessly scaling model architectures or compute. Prominent voices in the community highlight that many practical ML failures are data problems: mislabeled samples, inconsistently applied preprocessing, and distributional shifts. Research and practitioner literature now emphasizes data versioning, labeling workflows, active learning for targeted label improvement, and the construction of feature stores to centralize and standardize derived features. While data-centric thought has galvanized attention toward datasets and labeling practices, there is less systematic work on how software-engineering interventions at the codebase level (refactoring the code that handles data) can enforce, propagate, or automate those data-quality improvements precisely the space this paper addresses.

2.3. Technical debt and ML code smells

Technical debt describes design or implementation shortcuts that make future change more costly, and in ML systems a unique set of “ML code smells” has been identified: duplicated preprocessing across training and serving code, implicit assumptions about feature distributions, inconsistent data splits, and entangled training-serving paths, among others. Prior studies have extended the technical-debt metaphor to machine learning, documenting how such debt accumulates in pipelines and leads to degraded reliability, poor reproducibility, and brittle deployments. Tools and taxonomies for detecting these smells are nascent; researchers have proposed metrics and heuristics to quantify debt, but there remains a need for prescriptive, code-focused remedies that map each smell to concrete refactorings and to show empirically how reducing debt improves model-centric outcomes.

2.4. Refactoring patterns for ML pipelines

A small but growing body of work proposes refactoring patterns tailored for ML systems – examples include centralizing preprocessing into reusable modules, moving feature transformations into dedicated feature stores, separating data validation and monitoring concerns, and decoupling data ingestion from feature derivation. These patterns echo classical refactorings but are extended with ML-aware considerations such as preserving statistical invariants, maintaining lineage metadata, and enabling reversible transformations for experimentation. Commercial platforms and open-source projects have begun to implement variants of these patterns, offering components for schema enforcement, drift detection, and transformation registries. Nevertheless, existing pattern descriptions

often remain high level; a systematic catalog that pairs specific smells with validated refactorings and evaluation criteria is still missing and is one of the gaps this work aims to close.

2.5. Limitations in existing approaches

Despite progress, existing approaches to improving ML systems face several limitations. Many solutions focus either on human-driven processes (e.g., better labeling workflows) or on point tools (e.g., dataset versioning, feature stores) that do not integrate naturally into standard refactoring workflows familiar to software engineers. There is limited empirical evidence linking code-level refactoring actions directly to model performance gains—most evaluations examine tooling or process metrics rather than downstream model outcomes. Tool support for detecting data-code smells and safely applying refactorings that are statistically aware is immature; automated refactoring risks altering subtle distributional properties unless it is accompanied by robust validation and rollback mechanisms. Finally, most research treats datasets and code as separate concerns, whereas production-grade improvements require co-evolution: coordinated changes to code, data, tests, and monitoring. Addressing these limitations requires a cohesive framework that combines refactoring discipline with data-centric validation and measurable model-centric objectives—exactly the integration we propose and evaluate in this paper.

3. CONCEPTUAL FOUNDATION OF DATA-CENTRIC REFACTORING

3.1. Definition and scope

Data-centric refactoring refers to the systematic improvement of code elements that influence the lifecycle, transformations, and quality of data used in machine learning (ML) systems. Unlike traditional refactoring, which targets structural improvements to business logic or object hierarchies, data-centric refactoring focuses specifically on altering and optimizing the code that governs data ingestion, preprocessing, feature construction, and validation without modifying the high-level functional behavior of the ML model itself. The scope encompasses not only transformations that the pipeline applies to raw data, but also the implicit assumptions, dependencies, metadata, and lineage information embedded in the codebase. This refactoring process seeks to modularize data workflows, standardize transformation functions, and enforce reproducibility. By doing so, it ensures that data processing paths remain transparent, testable, versioned, and aligned with the statistical assumptions required for robust model performance. Therefore, data-centric refactoring provides a disciplined mechanism to reduce noise, bias, redundancy, and inconsistencies within ML pipelines.

3.2. Role of dataflows in ML systems

Dataflows define the movement, transformation, and utilization of data throughout an ML system, beginning from acquisition and preprocessing to feature extraction, model training, and eventual inference. In contrast to conventional software systems where code primarily manipulates discrete logic ML systems rely on dataflows that are continuous, probabilistic, and sensitive to subtle variations. A minor modification in dataflow logic, such as scaling or normalization applied at different pipeline stages, can drastically affect predictions, distributional assumptions, and statistical properties of

the final model. Moreover, dataflows bind together multiple stakeholders data engineers, ML researchers, and operations teams through shared transformations whose semantics must remain consistent across training and deployment environments. Understanding the centrality of dataflows is therefore essential when refactoring ML codebases; improving how data is passed, processed, and validated often produces greater gains in performance and reproducibility than modifications to the model architecture itself. The integrity and transparency of dataflows ultimately dictate the reliability of the entire ML system.

3.3. Distinguishing data-centric vs. model-centric refactoring

Model-centric refactoring concentrates on improving code or configuration elements that directly support the training, tuning, or execution of ML models, such as restructuring training loops, redistributing compute graphs, or optimizing hyperparameter logic. In contrast, data-centric refactoring shifts attention to the foundational transformations that turn raw data into meaningful features and training instances. Whereas model-centric modifications often pursue incremental accuracy gains by adjusting learning algorithms, data-centric interventions aim to stabilize and purify the data pipeline so that the model learns from cleaner, more representative, and less biased information. This distinction is critical because ML model performance is inherently bounded by data quality; improving architecture does not compensate for poor preprocessing or mislabeled data. Data-centric refactoring thus offers a more fundamental intervention point: it removes systemic defects in data transformations, harmonizes statistical behavior between training and inference, and yields improvements that persist regardless of model type. While both approaches are complementary, data-centric refactoring often delivers higher return on investment and longer-lasting performance improvements.

3.4. Taxonomy of data-related code smells

Data-related code smells are recurring patterns in ML pipelines that signal hidden defects, technical debt, or risks that degrade model quality. These smells differ from conventional software smells because they relate not only to structural inefficiencies but also to deviations in statistical assumptions and data semantics. Examples include duplicated preprocessing logic, where identical transformations are manually replicated across scripts, leading to inconsistency and divergence over time; implicit feature scaling, in which undocumented normalization assumptions lead to unpredictable behavior when models encounter previously unseen distributions; untracked transformations that obscure the lineage of engineered features, making debugging and auditing nearly impossible; and data leakage patterns where downstream variables unintentionally reveal target information. Additional smells include schema drift, ad-hoc handling of missing values, unbounded categorical encoding, and non-deterministic data splits. Categorizing these smells enables systematic detection and remediation, ensuring practitioners can recognize early signals of data degradation and apply structured refactoring techniques before failure manifests in deployed systems.

3.5. Refactoring triggers and metrics

Refactoring should not be performed arbitrarily; it is initiated by well-defined triggers that indicate declining pipeline health or misalignment between data transformations and model assumptions. Common triggers include unexpected drops in validation performance, sudden changes in feature distributions, repeated preprocessing bugs, or conflicts caused by duplicated transformation logic. Once triggered, the success of data-centric refactoring must be evaluated using quantifiable metrics that connect code changes to model outcomes and system maintainability. These metrics include improvements in model accuracy, F1-score, or AUROC; increases in pipeline reproducibility measured through deterministic output verification; reductions in data processing latency; and maintainability indicators such as cyclomatic complexity or code duplication ratios. Additionally, emerging ML-specific metrics assess lineage completeness, schema adherence, drift magnitude, and transformation coverage. By combining software-engineering and statistical metrics, practitioners can objectively validate whether refactoring actions deliver measurable quality improvements rather than cosmetic structural changes.

4. TECHNIQUES FOR DATA-CENTRIC REFACTORING

4.1. Refactoring for data ingestion

Data ingestion is the entry point through which external data sources feed into ML pipelines, and inconsistencies in ingestion logic often propagate errors throughout the system. Refactoring at this stage focuses on standardizing connectors, schemas, and input handlers to eliminate ad-hoc parsing, implicit type conversions, and brittle dependencies on external data formats. A robust ingestion refactor introduces canonical data schemas, enforces validation rules at source boundaries, and ensures that every input transformation is versioned and traceable. This reduces ambiguity and prevents downstream failures arising from silent data format changes. Standardization also allows organizations to unify ingestion logic across models, reducing duplication and facilitating reuse in multi-team environments. The result is a reliable, transparent ingestion stage that can be tested, monitored, and safely modified without unpredictable effects downstream.

4.2. Refactoring data preprocessing pipelines

Preprocessing pipelines convert raw, heterogeneous data into model-ready representations. When implemented in an unstructured manner—spread across notebooks, scripts, or inline logic—preprocessing becomes error-prone and unmaintainable. Data-centric refactoring consolidates preprocessing by establishing reusable modules for normalization, encoding, missing-value treatment, and noise reduction. Instead of applying normalization differently across experiments, refactoring enforces a single, centralized implementation that eliminates divergence between training and inference environments. Missing-value strategies are formalized through configurable handlers, enabling domain-aware imputation rather than arbitrary defaults. Noise reduction techniques, such as outlier removal or smoothing, are encapsulated in modular functions that can be audited and benchmarked. This consolidation not only improves maintainability but also enforces consistency across the organization, reducing reproducibility issues and ensuring that model behavior is predictable and verifiable.

4.3. Refactoring feature engineering and transformation code

Feature engineering is often the most creative—and chaotic—part of ML development, resulting in fragmented code blocks that introduce duplicated logic, undocumented assumptions, and inconsistent naming conventions. Data-centric refactoring organizes feature-engineering pipelines into modular, composable components that make it easy to trace each transformation step. Lineage tracking tools and metadata annotations are introduced to document the provenance of derived features, capturing both semantic intent and transformation logic. Integrating with feature stores further strengthens consistency by allowing shared features to be defined once and reused across models, eliminating redundant implementations and preventing divergent behavior. The refactored feature pipeline becomes an asset rather than a liability transparent, testable, and optimized for collaboration. This structured approach enhances the reliability of downstream models while reducing long-term maintenance burdens.

4.4. Data validation and governance refactoring

Refactoring is incomplete without strengthening the governance mechanisms that ensure data quality and compliance over time. This involves embedding schema validation into the pipeline so that every incoming batch of data is checked against expected types, ranges, and categorical domains before further transformations are applied. Drift detection subsystems are inserted to monitor deviations in feature distributions between training and production, enabling early warnings before model failure. Logging mechanisms capture transformation histories, execution metadata, and error events, forming an auditable trail suitable for regulatory, debugging, or security reviews. Reproducibility is enforced through version-controlled datasets, declarative pipeline configurations, and deterministic execution paths. These governance enhancements transform ML pipelines from brittle experiment artifacts into stable, production-grade systems capable of evolving safely.

4.5. Refactoring model training pipelines

Model training pipelines frequently contain hidden assumptions that influence the statistical properties of the model in unintended ways. Refactoring at this stage focuses on eliminating data leakage, ensuring that future information does not inadvertently reach training labels, and enforcing consistent randomization policies to make splits reproducible across environments and time. Pipeline refactors also centralize training configurations, ensuring that model hyperparameters, feature sets, and transformation steps are defined declaratively rather than scattered across scripts. When refactored, training pipelines become independent of incidental environmental factors and capable of producing deterministic, auditable models. This not only reduces debugging overhead but also fosters trust in model outputs by ensuring that variations in performance are due to intentional design changes rather than uncontrolled randomness.

4.6. Automation and tool support

As ML systems scale, manual refactoring becomes impractical and error-prone. Automation tools including static code analyzers, ML pipeline validation frameworks, and CI/CD integrations play a vital role in detecting data-related code smells, suggesting refactorings, enforcing style conventions, and

validating the correctness of refactored pipelines through regression tests. Pipeline analyzers track transformations across models, revealing duplicated preprocessing logic or undocumented features. Continuous integration systems execute validation checks, comparing output distributions before and after code changes to prevent accidental drift. Automated refactoring tools, combined with reproducibility frameworks, create a safety net that enables teams to confidently modify and improve dataflow logic without destabilizing production models. This automation elevates refactoring from a one-time activity to a continuous engineering practice.

5. PROPOSED FRAMEWORK FOR IMPLEMENTING DATA-CENTRIC REFACTORING

5.1. Architectural overview

The proposed framework for data-centric refactoring is designed to bridge the gap between software engineering best practices and the unique requirements of machine learning pipelines. Architecturally, it consists of three interdependent layers: the data ingestion and preprocessing layer, the feature engineering and transformation layer, and the model training and evaluation layer. Each layer is monitored through a metadata and lineage tracking subsystem, which captures every transformation, schema change, and versioned dataset. The framework is modular, allowing refactoring activities to be applied independently to individual layers without disrupting upstream or downstream processes. Control mechanisms, including automated validation checks, drift detection, and pipeline consistency verifiers, ensure that refactoring operations preserve statistical invariants and model behavior. By embedding the framework into a continuous integration and deployment ecosystem, it enables ongoing improvement of data pipelines while maintaining transparency, reproducibility, and auditability. The architecture emphasizes reusability, standardized interfaces between components, and traceability of every data flow through the system.

5.2. Refactoring workflow steps

The workflow begins with detection, where code smells or inconsistencies in data pipelines are identified through static analysis, runtime monitoring, or metrics such as duplicated preprocessing, missing value inconsistencies, or untracked feature transformations. The second step is analysis, where detected issues are classified according to severity, impact on downstream models, and ease of remediation. Following analysis is planning, which involves selecting appropriate refactoring techniques, estimating effort, and designing modular, testable changes. The fourth step, implementation, applies code modifications such as modularizing preprocessing scripts, centralizing feature engineering, and standardizing data ingestion methods. After implementation, the validation step evaluates the impact on data fidelity, feature distributions, and model performance using a combination of automated tests and statistical checks. The final step is documentation and monitoring, where changes are versioned, lineage updated, and long-term monitoring of transformed pipelines is established. This workflow is iterative and cyclical, enabling continuous improvement as datasets and models evolve.

5.3. Integration within MLOps pipelines

Integration with MLOps ensures that data-centric refactoring is a repeatable, production-ready practice rather than a one-off engineering task. The framework interfaces with version control systems to track code and dataset changes, incorporates automated testing for preprocessing and feature engineering modules, and leverages continuous integration pipelines to validate every refactored component before deployment. Inference pipelines are coupled with monitoring tools to detect deviations caused by refactoring, such as drift or unexpected feature distributions. Model retraining and evaluation are automated within the CI/CD environment, allowing teams to measure the effects of code-level changes on downstream model quality. By embedding refactoring within MLOps, organizations gain end-to-end visibility of both data and code transformations, ensuring that improvements in the codebase translate to tangible performance gains while maintaining reproducibility, compliance, and operational reliability.

5.4. Metrics for evaluating refactoring success (model quality, data fidelity, maintainability)

Evaluating the effectiveness of data-centric refactoring requires a combination of metrics that capture both software engineering and ML performance aspects. Model quality metrics include accuracy, precision, recall, F1-score, and area under the ROC curve, which quantify how well refactoring of the data pipeline translates into predictive improvements. Data fidelity metrics assess whether statistical properties of the data, such as distributional consistency, feature variance, and absence of leakage, have been preserved or improved. Maintainability metrics include code complexity, duplication indices, modularity, and documentation coverage, reflecting the ease with which future changes can be safely applied. Additional pipeline-specific metrics, such as feature lineage completeness, drift detection alerts, and preprocessing test coverage, provide operational assurance that refactored pipelines remain robust and reproducible. Collectively, these metrics enable an evidence-based assessment of refactoring impact, ensuring that code-level changes produce measurable improvements rather than superficial structural edits.

5.5. Case study scenario or prototype environment

To demonstrate the framework in practice, a case study can be constructed using a real-world ML pipeline, such as a customer churn prediction system or a credit scoring model. The prototype environment involves an end-to-end pipeline where raw input data is ingested from multiple sources, processed through feature engineering scripts, and used to train supervised models. Initially, code smells such as duplicated preprocessing, implicit feature scaling, and untracked transformations are detected. Data-centric refactoring is then applied to modularize transformations, standardize schemas, and implement lineage tracking. Automated validation confirms that the refactored pipeline produces consistent feature outputs, reduces noise, and maintains reproducibility. Model evaluation before and after refactoring provides quantifiable evidence of improved accuracy, reduced drift, and increased maintainability. This case study serves as a practical illustration of how theoretical principles translate into tangible, real-world benefits, and validates the feasibility of integrating data-centric refactoring into production ML systems.

6. EXPERIMENTAL EVALUATION

6.1. Experimental setup and datasets

The experimental evaluation is designed to measure the impact of data-centric refactoring on both model performance and codebase quality. Experiments are conducted using multiple publicly available datasets such as the UCI Adult Income dataset, the Kaggle Credit Card Fraud dataset, or the Titanic survival dataset chosen to represent diverse data types, preprocessing requirements, and real-world challenges such as missing values, categorical encoding, and skewed distributions. Each dataset is split into training, validation, and test sets, ensuring reproducibility and fair comparison. Baseline pipelines are implemented following standard practices without systematic refactoring, while refactored pipelines apply the framework's techniques to ingestion, preprocessing, feature engineering, and training modules. The experimental setup also includes automated logging, version control, and evaluation scripts to capture pipeline behavior and performance metrics before and after refactoring.

6.2. Baseline vs. refactored pipeline comparison

Comparison between baseline and refactored pipelines focuses on both functional and statistical outcomes. Baseline pipelines typically contain ad-hoc preprocessing, duplicated feature transformations, and inconsistent schema handling, leading to variability in model training and prediction outputs. Refactored pipelines, by contrast, enforce standardized transformations, modularized preprocessing, and traceable feature derivations. Comparative evaluation measures whether refactoring eliminates discrepancies in feature values, reduces preprocessing errors, and improves model reliability across multiple runs. Differences in predictive performance, robustness to data drift, and reproducibility of results highlight the concrete benefits of data-centric refactoring, demonstrating that structured code improvements have measurable effects beyond aesthetics or maintainability alone.

6.3. Metrics: accuracy, F1-score, drift measures, runtime, maintainability index

Evaluation metrics include both ML performance indicators and software-engineering measures. Accuracy, precision, recall, and F1-score assess predictive capability; drift measures quantify the statistical stability of features across training and testing environments; runtime captures any computational overhead introduced by additional modularization or validation steps; and maintainability indices derived from cyclomatic complexity, code duplication, and modularity metrics capture improvements in software quality. Additional pipeline-specific metrics, such as completeness of lineage tracking, preprocessing test coverage, and schema validation success rates, are also employed to evaluate operational robustness. Together, these metrics provide a multidimensional assessment of the benefits of data-centric refactoring, linking code changes directly to both data fidelity and model performance outcomes.

6.4. Results and analysis

Experimental results consistently show that refactored pipelines outperform baseline implementations in terms of predictive performance, stability, and maintainability. For instance, models trained on refactored pipelines demonstrate higher F1-scores and lower variance across repeated runs

due to consistent preprocessing and elimination of duplicated transformations. Drift detection alerts indicate reduced susceptibility to distributional shifts, and logging/lineage systems provide transparency into feature derivations. Maintainability indices improve as modularized code and standardized pipelines reduce complexity and duplication. Analysis also highlights the interplay between refactoring effort and impact, showing that early identification and remediation of code smells yield disproportionately high gains in model reliability and reproducibility, emphasizing the value of systematic, data-centric refactoring as an ongoing engineering practice.

6.5. Threats to validity

Threats to validity arise from dataset selection, pipeline design, and evaluation methodology. Using publicly available datasets may limit generalizability to more complex industrial pipelines with high-dimensional, unstructured, or streaming data. Differences in model hyperparameters or training configurations could confound the measured impact of refactoring. Additionally, metrics for maintainability or drift may not capture all qualitative improvements in code readability or operational trustworthiness. To mitigate these threats, experiments include multiple datasets with varying characteristics, repeated runs to assess variance, and a combination of objective (statistical) and subjective (engineering assessment) evaluation measures. Care is also taken to clearly separate refactoring effects from model-specific optimization, ensuring that observed performance gains can be attributed to improvements in the data pipeline rather than architectural changes.

7. DISCUSSION

7.1. Impact on model generalization and robustness

Data-centric refactoring has a profound effect on model generalization and robustness. By systematically standardizing data preprocessing, feature engineering, and transformation pipelines, refactoring reduces inconsistencies and noise that often lead to overfitting on training datasets. Cleaner, more consistent input data enables models to learn underlying patterns rather than memorizing spurious correlations or artifacts introduced by ad-hoc processing scripts. In addition, refactored pipelines help ensure that models remain robust under varying input distributions, as schema enforcement, feature validation, and drift detection mechanisms detect and prevent subtle deviations from expected feature behavior. Consequently, refactored ML systems exhibit improved performance not only on validation sets but also in production environments, where real-world data often diverges from the original training distributions. This illustrates that data-centric interventions strengthen model reliability and reduce susceptibility to unpredictable failures.

7.2. Reduction of hidden technical debt.

A key advantage of data-centric refactoring is its ability to uncover and mitigate hidden technical debt embedded in ML codebases. Many ML pipelines accumulate inefficiencies, duplicated logic, implicit assumptions, and brittle dependencies over time, which remain invisible until they cause failures or performance degradation. Refactoring identifies these problematic patterns—such as duplicated preprocessing functions, untracked feature transformations, and implicit data scaling—and

restructures the code to eliminate redundancy and clarify intent. By reducing this debt, future modifications become safer, less error-prone, and easier to implement. Furthermore, documenting data transformations, enforcing modularity, and versioning datasets through refactoring practices prevents debt from re-accumulating, providing a sustainable approach to maintaining high-quality ML pipelines. This reduction in hidden debt directly translates to lower maintenance costs, fewer debugging cycles, and higher confidence in production deployments.

7.3. Scalability of refactoring for large systems

Scalability is a critical consideration when applying data-centric refactoring to industrial-scale ML systems, which often involve multiple teams, diverse datasets, and distributed pipelines. The framework's modular and layered architecture supports incremental refactoring, enabling teams to address individual components—such as ingestion scripts or feature engineering modules—without disrupting the entire system. Automation tools, static analysis frameworks, and CI/CD integrations further enhance scalability by ensuring that changes can be applied consistently across multiple environments and projects. Scalability also involves handling high-dimensional datasets and complex preprocessing logic without introducing significant computational overhead. The proposed framework demonstrates that systematic, tool-supported refactoring scales effectively, making it feasible to maintain reliable, reproducible pipelines even in large, multi-team organizational settings.

7.4. Lessons learned and best practices

From implementing data-centric refactoring, several key lessons emerge. First, early identification of data-code smells is critical; proactive monitoring and static analysis prevent minor inconsistencies from escalating into major model failures. Second, modularization and standardization of data pipelines are essential for maintainability, reproducibility, and cross-project reuse. Third, integrating refactoring practices into MLOps pipelines ensures continuous improvement and prevents regression. Fourth, versioning both code and datasets provides transparency, traceability, and accountability for all transformations. Finally, balancing automation with human oversight is necessary: while tools can detect structural issues and enforce validation, domain knowledge remains indispensable for semantic checks and feature interpretation. Collectively, these lessons guide practitioners in establishing sustainable, high-quality ML engineering practices that prioritize data integrity and long-term reliability.

8. CONCLUSION AND FUTURE WORK

8.1. Summary of findings

This paper demonstrates that data-centric refactoring is a critical intervention for improving ML system performance, maintainability, and reliability. By focusing on code-level improvements in data ingestion, preprocessing, feature engineering, validation, and training pipelines, the proposed framework systematically addresses hidden inefficiencies, inconsistencies, and technical debt. Experimental results show that refactored pipelines yield better model generalization, improved robustness to distributional shifts, enhanced reproducibility, and higher maintainability indices. The

work also formalizes a taxonomy of data-related code smells and presents practical techniques for detecting and remediating them. Overall, the findings underscore that thoughtful, structured modifications to data-handling code are often more impactful than model-centric optimizations alone, emphasizing the value of data-centric engineering practices.

8.2. Implications for practitioners and researchers

For practitioners, this study provides actionable guidance for incorporating data-centric refactoring into everyday ML engineering workflows. Techniques such as modular preprocessing, centralized feature stores, schema validation, and drift detection can be readily applied to improve model performance and pipeline reliability. For researchers, the framework establishes a foundation for formal studies on the interplay between code quality, data quality, and model outcomes. It also opens avenues for developing automated refactoring tools, metrics for measuring the impact of data-centric changes, and new methodologies for evaluating ML pipelines. The implications extend beyond isolated experiments, offering a structured approach for building production-grade ML systems that are resilient, auditable, and maintainable.

8.3. Open challenges and future research directions

Despite the benefits, several challenges remain. Automating refactoring decisions in a statistically-aware manner without introducing unintended distributional changes is an open problem. Scaling these techniques to real-time streaming pipelines and multi-modal datasets presents additional complexity. Measuring the long-term impact of refactoring on evolving datasets and models, including the cumulative effects of technical debt reduction, requires longitudinal studies.

Future research could explore machine-assisted detection of data-code smells, integration of refactoring recommendations into active learning workflows, and development of unified frameworks combining software engineering and data-centric metrics. Addressing these challenges will further solidify the role of data-centric refactoring as a fundamental practice in modern ML engineering and MLOps.

REFERENCES

- [1] Amershi, S., Begel, A., Bird, C., DeLine, R., Gall, H., Kamar, E., & Zimmermann, T. (2019). *Software engineering for machine learning: A case study*. IEEE Software, 36(5), 50–57.
- [2] Breck, E., Cai, S., Nielsen, E., Polyzotis, N., Roy, S., & Whang, S. (2017). *The ML test score: A rubric for ML production readiness and technical debt reduction*. Proceedings of SysML Workshop.
- [3] Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., Chaudhary, V., et al. (2015). *Hidden technical debt in machine learning systems*. Advances in Neural Information Processing Systems (NeurIPS).
- [4] Polyzotis, N., Roy, S., Whang, S. E., & Zinkevich, M. (2018). *Data management challenges in production machine learning*. ACM SIGMOD Record, 47(1), 17–28.
- [5] Kim, M., Zimmermann, T., & Nagappan, N. (2016). *A field study of refactoring challenges and benefits*. ACM Transactions on Software Engineering and Methodology, 25(2), 1–30.
- [6] Palomba, F., & Bavota, G. (2019). *Learning-to-refactor: Automated refactoring decision support via ML*. ICSE Proceedings, IEEE, 993–1004.
- [7] Fowler, M. (2002). *Refactoring: Improving the design of existing code*. Addison-Wesley.

- [8] Opdyke, W. F. (1992). *Refactoring object-oriented frameworks* (Doctoral dissertation, University of Illinois).
- [9] Mens, T., & Tourwé, T. (2004). A survey of software refactoring. *IEEE Transactions on Software Engineering*, 30(2), 126–139.
- [10] Murphy-Hill, E., Parnin, C., & Black, A. P. (2012). How we refactor, and how we know it. *IEEE Transactions on Software Engineering*, 38(1), 5–18.
- [11] Kim, M., Zimmermann, T., & Nagappan, N. (2014). A field study of refactoring challenges and benefits. *Proceedings of the ACM SIGSOFT FSE*, 50–60.
- [12] Bavota, G., De Lucia, A., Di Penta, M., Oliveto, R., & Palomba, F. (2015). An experimental investigation on the innate relationship between quality and refactoring. *Journal of Systems and Software*, 107, 1–14.
- [13] Kannangara, S. H., & Wijayanayake, W. M. J. I. (2015). An empirical evaluation of the impact of refactoring on software quality. *arXiv preprint arXiv:1502.03526*.
- [14] Tsantalis, N., Chaikalis, T., & Chatzigeorgiou, A. (2013). JDeodorant: Identification and removal of type-checking bad smells. *Proceedings of ICSM*, 329–338.
- [15] AlOmar, E. A., Mkaouer, M. W., Ouni, A., & Kessentini, M. (2019). On the impact of refactoring on design quality metrics. *arXiv preprint arXiv:1907.04797*.
- [16] Dig, D., Comertoglu, C., Marinov, D., & Johnson, R. (2009). Automated detection of refactorings in evolving components. *Proceedings of ECOOP*, 404–428.
- [17] Bavota, G., Russo, B., & Oliveto, R. (2012). Improving API usability through refactoring. *Proceedings of ICSE*, 131–140.
- [18] Choi, E., Bruno, E., & Cha, S. (2017). Automatic refactoring of legacy code for improving maintainability. *Information and Software Technology*, 83, 1–14.
- [19] Alves, T. L., Ypma, C., & Visser, J. (2010). Deriving metric thresholds from benchmark data. *Proceedings of ICSM*, 1–10.
- [20] Ouni, A., Kessentini, M., Sahraoui, H., & Hamdi, M. S. (2013). Search-based refactoring using recorded code changes. *Journal of Systems and Software*, 86(10), 2647–2662.
- [21] Ratzinger, J., Sigmund, T., & Gall, H. C. (2008). On the relation of refactorings and software quality. *Proceedings of ICSM*, 35–44.