



Cloud-Native Agentic AI for Insurance Fraud Detection

Michael Anderson

Asst Professor

Abstract

The study investigates DevOps-driven data engineering for agentic AI in insurance fraud detection with a focus on cloud-native automations. Data are centrally secured, accessible, governed, of reproducible quality, properly profiled, and efficiently prepared. Design, development, and deployment encompass extensive safeguards for risk-sensitive decision-making. A superset of continuous integration/continuous deployment tailored to machine learning as well as data governance underlie the approach, demystifying the technology stack and enabling success at operational scale. The investigation is motivated by regulatory, economic, and ethical pressures on the industry, demand for predictive solutions, and pain points of legacy monitoring systems. Automation of machine-learning model development and deployment mitigates issues of scalability, reproducibility, and security while incorporating risk assessment.

Organizations worldwide are increasingly turning to artificial intelligence in a bid to detect fraudulent activity in all its forms. Such systems, however, generally stem from one-off proof-of-concept initiatives, operate in isolation, and lack sufficient controls for operational deployment in support of everyday decision-making. Nevertheless, agentic artificial intelligence—actant solutions capable of executing decisions without direct human intervention—promises a transformative competitive advantage for business leaders. Integrating governance and control ensures that these automated systems operate within an acceptable risk tolerance, align with organizational values, and minimize unintended consequences. Insurance fraud detection solutions typically fall short of these requirements.

Keywords: DevOps Driven Data Engineering, Agentic Artificial Intelligence, Insurance Fraud Detection, Cloud Native Automation, Machine Learning Operations, Continuous Integration Continuous Deployment, Data Governance Frameworks, Risk Sensitive Decision Making, Automated Model Deployment, Predictive Fraud Analytics, Scalable AI Systems, Reproducible Data Pipelines, Secure Data Architectures, Regulatory Compliance In Insurance, Ethical AI Governance, Legacy System Modernization, Autonomous Decision Systems, Operational AI Control, Model Lifecycle Management, Enterprise AI Adoption.

1. Introduction

Insurance fraud, encompassing deliberate misreporting of claims, is a costly global problem that erodes corporate earnings and increases annual customer expenses. The prevention of fraudulent activity requires constant vigilance by insurance companies, which dedicate substantial resources to fraud detection. However, the sheer volume of claims makes manual assessment prohibitively expensive. Consequently, organizations deploy analytical models capable of detecting suspicious and potentially fraudulent claims for further investigation—ideally using enterprise

data lakes that integrate multiple sources of structured and unstructured information and external databases. Machine Learning (ML) processes are well-suited to this task, as they can analyze the information within, derive predictive patterns, and then automatically apply those patterns to future claims.

Nevertheless, the application of ML models to support fraud detection remains constrained by a range of factors. Data and ML development remain mainly siloed, thus preventing stakeholder-driven, responsive development, while produce-scand-operate processes remain disjointed

rather than collaborative and continuous. The resulting tools lack transparent governance and control, leading to security concerns and limiting their deployment in production systems; the empirical evidence for the payoff from such AI still remains patchy. Hence, the models cannot be deployed as services, scaling the analytic capability of the enterprise, or rapidly operationalized via pipelines that enable continuous integration/continuous delivery (CI/CD) through data.

1.1. Objectives and Scope of the Study

The primary goal of this study is to explore, develop, and validate a cloud-native DevOps-automated deep learning system to detect fraudulent insurance claims. While the agentic AI perspective offers the promise of scaling machine learning (ML) model training, monitoring, and evaluation, the principal emphasis here is on how the combination of DevOps and certain elements of agentic AI can bring transparency, security, governance, and auditability to the data preparation and model development life cycles. Additional issues under the DevOps umbrella—reproducibility, continuous integration/continuous deployment (CI/CD), supply-chain security, and quality—are also significant but less critical. Nevertheless, testing and evaluation across nearly all these dimensions will maximize usefulness and reliability.

Research activity in the insurance sector has long focused on fraud, reflecting the extensive financial losses incurred around the world. Despite the gradual implementation of machine learning solutions along the fraud-detection process, the application of state-of-the-art (SOTA) deep learning models remains limited. The main drivers of this gap are the complexity of deep learning networks and the difficulties involved in data preparation. Nevertheless, interest in insurance fraud detection persists, and recent developments in data engineering enabled by a combination of cloud-native DevOps and SOTA models can enhance scalability, reproducibility, and security while providing the infrastructure required for the deployment of agentic AI. Such developments could therefore increase the attractiveness of deep learning, closing the gap between established and emerging methods.

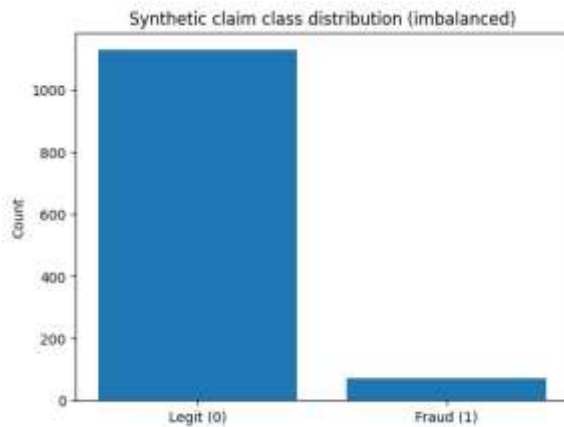


Fig 1: Scaling Deep Learning for Fraud Detection: A Cloud-Native DevOps Framework for Automated Governance and Agentic AI Integration

2. Background and Motivation

Detecting fraud during insurance transactions has grown into a pressing concern for the industry. Reports from several countries confirm the growth of fraudulent transactions, together with a lack of transparency of the underlying data. Countries such as the Netherlands have introduced regulatory measures to tackle these concerns; such measures also increase pressure on the insurance industry to invest in machine learning solutions that support transparency and compliance while improving detection performance. While recent years have seen a substantial rise in interest in the research community, associations and industry, most efforts are still in a proof-of-concept stage – few, if any, results have made it into production. The growing importance of fraud detection as part of the insurance value chain creates additional pressures:

developing models that work at scale using industry-grade machine learning practices. Combining the principles established by DevOps for Data Engineering with industry-graded development and deployment considerations allows the creation of solutions that address these drivers. Scalability of the developed and integrated solutions appears to be the term that has been leant on most heavily; fulfilling the additional requirements of compliance, governance, transparency and reproducibility is a clear by-product of the DevOps-for-Data-Engineering principles. Beyond this, it is expected that the underlying business problems in the case of insurance transaction fraud detection can be sufficiently addressed; using the performance of the models in the support of insurance companies and associations can be a key source of motivation.



Equation 1: Data preparation equations

1.1 Min–Max normalization (MinMaxScaler)

Let:

- original feature value = x
- minimum of the feature column = $x_{\min}$
- maximum of the feature column = $x_{\max}$

Goal: map $x \in [x_{\min}, x_{\max}]$ to $\hat{x} \in [0,1]$.

Step-by-step derivation

- Shift so the minimum becomes 0

$$x - x_{\min}$$

- Scale so the maximum becomes 1
The max shifted value is $x_{\max} - x_{\min}$. Divide by that:

$$\hat{x} = \frac{x - x_{\min}}{x_{\max} - x_{\min}}$$

- (General range $[a, b]$ if needed)
First compute $\hat{x} \in [0,1]$ as above, then stretch/shift:

$$x' = a + (b - a)\hat{x} = a + (b - a) \frac{x - x_{\min}}{x_{\max} - x_{\min}}$$

2.1. Key Drivers of Research and Innovation

Scaling, reproducibility, security, and governance are becoming increasingly relevant for deploying ML solutions in production. The insurance sector is facing continuous pressure to deliver models that address key business questions and support operational decisions. However, the industry has yet to embrace innovation at scale, often relying on traditional technologies, bespoke approaches, and partly redundant ML initiatives that are rarely integrated into business processes. Pressure for innovation is mounting in insurance companies and new insurance players are entering the market, forcing the insurance industry to compete through continued innovation. Agentic AI comprises continuous integration and delivery principles applied to ML workflows, enabling the creation of automated, monitored pipelines to regularly produce and deploy new ML models. Regulatory requirements such as the General Data Protection Regulation are forcing insurance companies to implement stricter data governance processes. Data ownership, access management, data quality, data classification, and data lineage tooling are becoming core data engineering issues. Insurance companies can leverage modernization and cloud services to integrate security into the design and transpose cloud-

native principles and patterns into data engineering and ML process automation.

Class	Meaning	Count	Share
0	Legitimate Claims	1129	94.1%
1	Fraudulent Claims	71	5.9%

3. Theoretical Framework

DevOps principles provide valuable underpinnings for scalable, reproducible data engineering. Agentic AI strives for intelligent self-management of complex systems. Together, these ideas structure the research and guide the design of cloud-native, deep-learning-enabled software that addresses longstanding pain points in insurance fraud detection.

Continuous integration/continuous deployment (CI/CD) for machine learning, data governance, automated monitoring, and ethical considerations shape the investigation. Cloud-native processing meshes streaming and batch data, engenders a data-automation layer that supports ingestion, access, security, quality, lineage, profiling, and transformation, and provides a second automation layer for cloud-native ML pipelines.

Deep-learning-automation-as-a-service solutions ease proof-of-concept work by growing the four Tenet-specific pains and supporting three more—lack of production-ready ML, reproducibility, and security—before being taken up by the business-facing ML crew. Service teams outside the O&M squad own related service performance, stability, and security at service-data pipelines that form the dumb enablers (for safety, et al.).

CI/CD for ML and agentic AI's Monitoring tenets telescope to deliver proof of concept, O&M grows support for six company Operations pain points, and the master plan then steers operationalization of automated FaaS and ML pipelines. The latter deal with workloads Dw and Dd—crafting, hosting, and automating Thessala, probably doing the same for any approved MLOps job-lot as proof of

concept, and supporting by touch-related-en)))) and Terraform/Docker-for-Dummies transports.



Fig 2: Synergizing DevOps and Agentic AI: A Cloud-Native Architectural Framework for Scalable Deep Learning and Automated Fraud Mitigation

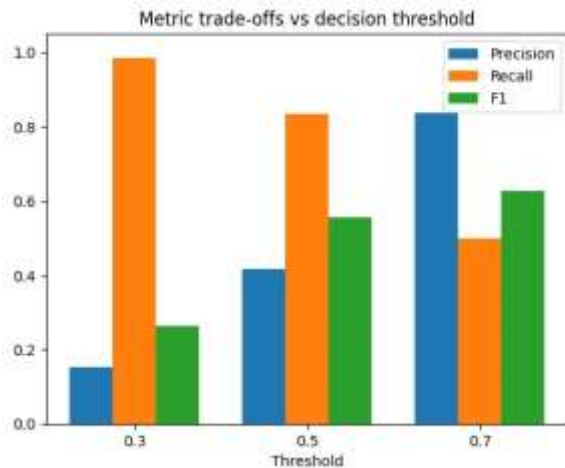
3.1. Conceptual Overview of Key Principles

Continuous Integration/Continuous Deployment of Machine Learning

Although DevOps was conceived for software development, its principles of continuous integration/continuous deployment (CI/CD) can also be applied to machine learning (ML) models to enable the timely availability of high-quality algorithms. CI/CD focuses on three distinct phases: a) continuous integration, which includes the development, testing, and validation of ML pipelines; b) continuous deployment of a candidate model that is updated regularly; and c) continuous training, which acknowledges that the underlying conditions affect the performance of deployed models. This requires an end-to-end ML pipeline (e.g. data ingestion, cleaning, preparation, modeling) to be in place, which, besides automating the modeling process, allows the management of model accuracy and robustness throughout the entire model development and production lifecycle.

Data Governance

Data are the lifeblood of successful ML solutions, and specific governance strategies, tailored to the scale and complexity of the organization, are needed to ensure that data is always available, secure, and compliant with regulatory requirements. Data governance comprises the design and execution of policies, standards, and processes that enable the acquisition and use of data, guarantee confidentiality, availability, and integrity, and mitigate relevant risks, such as security breaches, data leaks, and regulatory noncompliance. Effective data governance provides an inventory of the data assets of an enterprise, including metadata such as data classification, schema, access rights, lineage, source, quality metrics, documentation, and contact information, to ensure data consistency, security, and quality. It creates a shared understanding of data assets, minimizes risks, and generates a cycle of ongoing improvement.



4. System Architecture

The architecture provides a high-level view of the solution, detailing components, data flows, and integration patterns. Each process involved in insurance fraud detection is framed as a pipeline. The choice for cloud-native deployment is justified, and the multi-layer automation

keeps complexity in check and security considerations addressed.

A. Data Ingestion Pipelines

Data ingestion covers the collection, access, lineage tracking, profiling, and quality assurance of data in separate processes for batch and streaming feeds. Access is controlled for both data producers and consumers according to governance principles. For batch loads, point-in-time access is provided (with snapshots enabled) and correctness ensured with integrity checks. A quality pipeline checks for records exceeding profile quality thresholds. Data sources include client lists (ground truth), transactional databases, claims history, user sign-up profiles, and social media postings.

B. Cloud-Native ML Pipelines

ML training is also organized as a set of pipelines. It comprises readiness checks, orchestration, containerization, training, monitoring, model registry, and CI/CD elements and supports the online production of models for direct consumption by the business. Readiness checks verify the secure setup of cloud resources, design specifications, and data availability. The orchestration component handles resource provisioning and model training. In terms of model reproducibility, a container image is prepared before training—containing the environment libraries, the model training code with hyperparameter definitions, and the model storage bucket. The scan module periodically verifies the training setup based on CI principles.

Thres hold	T P	F P	F N	T N	Preci sion	Rec all (TP R)	F1 Sc ore	FP R
0.30	7 0	3 9 2	1	73 7	0.15	0.9 9	0.2 6	0. 35
0.50	5 9	8 3	1 2	10 46	0.42	0.8 3	0.5 6	0. 07
0.70	3 6	7	3 5	11 22	0.84	0.5 1	0.6 3	0. 01

4.1. Data Ingestion and Governance

Ingestion pipelines ensure quality, governance, and security, delivering trustworthy datasets for consumption. A Batch Loading Zone (BLZ) maintains structured data accessible for business users and analysts. The Cleaning Data Repository (CDR) holds refined datasets ready for machine learning. Stream and batch ingestions differ for incremental fraud monitoring and periodic training data updates. Batch operations handle files staged in the Cloud Storage Landing Zone. A comprehensive audit trail, established via DataFlow, tracks every operational aspect. Automated data quality checks, covering schema validation, completeness, accuracy, timeliness, and consistency, are vital for archived fraud detection data. For sensitive data, tools like Double-Check create a secondary copy with schema masking, ensuring secure development/testing access. Governance mechanisms enable appropriate authorizations at every stage. Data egress employs Data Transfer Service with fpp-level control, and Data-Lake-Information-Climate exposes additional technical metadata (e.g., data quality scores).

Equation 2: Outlier capping at three standard deviations

Let:

- dataset values: $x_1, \dots, x_n$

- mean:

$$\mu = \frac{1}{n} \sum_{i=1}^n x_i$$

- population standard deviation:

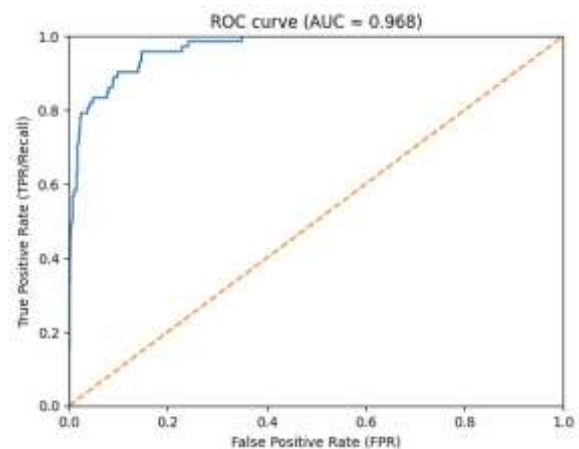
$$\sigma = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2}$$

- bounds:

$$L = \mu - 3\sigma, \quad U = \mu + 3\sigma$$

Capping rule (piecewise)

$$x^{(cap)} = \begin{cases} L, & x < L \\ x, & L \leq x \leq U \\ U, & x > U \end{cases}$$



4.2. Cloud-Native Machine Learning Pipelines

Pipelines ingesting the data into the machine learning (ML) module enable both model training and model orchestration at inference time. ML candidate models are deployed to a Google Cloud Vertex AI model registry where they can be evaluated, versioned, and monitored. Containerization enables new model versions to be generated in CI/CD for ML style, leveraging the full power of cloud-native continuous integration pipelines. New models can be trained with new hyperparameter searches when new training data is available, or when a new candidate architecture needs to be explored from among a shortlist. As the training process is encapsulated in a container, failure to train, timeout, or over-consumption of resources are easily preventable within limits automatically enforced by the orchestration system. Quality of training and production candidate models is enabled by automating the processes of building datasets, data and model profiling, exploring trained models, and discovering and evaluating hyperparameter search spaces. Automated pipelines of operational characteristics indicated previously for data ingestion can also be used to prepare datasets ready for training or retraining of a fraud detection setup. Careful control of the relational algebra

underlying these preparations underpins the constant availability of datasets suitable for high quality reproduction of fraud detection candidate models. A separate process regularly evaluates chessboard comparison matrices of candidate fraud detection model performance on representative datasets. Such processes help to remove unsuitable candidate models from the registry before reuse at inference time, augmenting the checks that should be in place by any ML Ops module. These candidate models are selected primarily by F1 score as precision and recall are commonly known to have traditionally opposing forces in the fraud domain.

5. Data Preparation and Feature Engineering

Extensive data preparation is necessary to transform in-house and publicly available datasets into appropriate model training artifacts. Apart from data cleaning and the normalization of categorical and continuous features, privacy-sensitive personal identifiers such as names, dates of birth, and addresses are removed. Normalization of sensitive features indirectly strengthens the governance around models exposed via APIs. Data from the skewed distribution of fraudulent cases must be selected carefully. Thus, synthetic samples that resemble the fraud detection process of insurance claim requests are added, and noise is injected into normal claims, enhancing robustness to overfitting. While the model considers simple features such as claim amount and policy number in its decisions, the Rewards and Fraud scenarios discussed earlier can signal the occurrence of higher fraud propensity in a broader set of features with disparate scales.

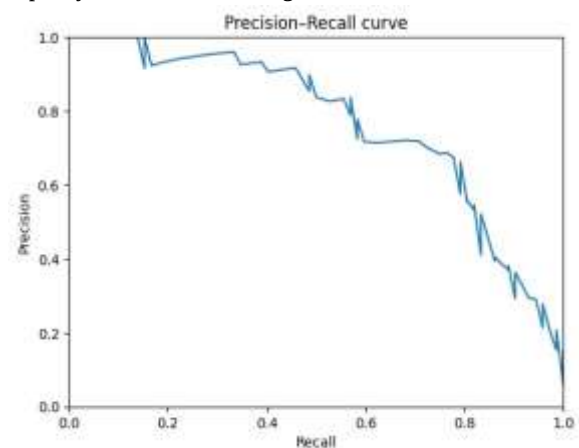
Data cleaning begins with the removal of duplicates, previously expelled claims from the combination of a unique policy number and claim amount, and all-zero feature vectors. The long and lat categorical features are then transformed into geographic location categorical features since they do not hold any semantic information for the detection process. Afterward, time dependency for key features is also injected. Finally, one-hot encoding is

employed for categorical features with more than thirty unique values. Normalization is subsequently performed using MinMaxScaler, which scales the values to a specified range (defaulting to 0, 1) while preserving base intervals. For other features with minor pI, outliers are capped, i.e., values beyond three standard deviations are replaced with the closest in-bound value.

Feature selection narrows down from thirty potential features to the following five indicators of insurance fraud propensity: digits in policy number, elapsed days since policy was taken, elapsed days since claim was lodged, last claim amount made on the policy, and time since last claim. Resampling is then performed to address class imbalance challenges during model training. The min-max scaling approach described earlier is used, injecting noise via the following function:

```
main.9933nelvalbitinewtjtgstqitunlirumrUfatitisww  
rufatittqwsjtbittwritnflutnmRflutraturwRflutditt  
nflutnitutatwj5wd9ttsrgklZflutnitnitwgl
```

ᐅᐅᐅᐅᐅᐅᐅ. Finally, after the garbage in–garbage out principle prompts the inclusion of multiple feature profiles, feature profiling and characterization across the entire data preparation process ensure data quality and forecasting capacity for machine learning models.



5.1. Data Cleaning and Transformation Techniques

The data preparation workflow relies on a combination of data-catalog services, ETL orchestrators, and user-



developed scripts or notebooks. The first step focuses on the cleansing of the InsuranceCompany Fraud Detection Dataset (ICFDD) (Kaggle 2021) prior to being used in the training, testing, and validation of machine learning models. After almost 1900 records of missing values were removed, the categorical attributes (underwriting, route, make, type, and fraud) had their values transformed into one-hot variables to allow for proper processing by machine learning algorithms. The fraud value was inversely coded via bitwise not and then multiplied by -1 to prepare the dataset for plotting of the Fraud Detection Fractal, thus producing a stronger contrast between fraud and nonfraud patterns. Continuous, discrete, and categorical attributes were then analysed using K-means profiling in order to leverage the data's natural structure for future model training.

A separate data preparation workflow conforms the ICFDD for ingestion into an NLP-based fraud detection pipeline. User-defined Python scripts perform additional cleaning duties, including removal of URL and HTML coding information that could otherwise interfere with the NLP engine's predictive capability. Normalization transforms any anomalous or redundant strings into a common format and part-of-speech tagging supplies structure to every word token in a description to facilitate intelligent selection of target language for the NLP engine. Word-count detection further identifies portions of the comments made by an insurance agent that may be inadequate for high-quality translation into the report's target language.

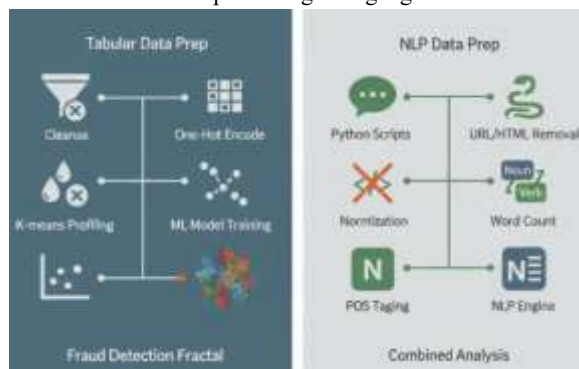


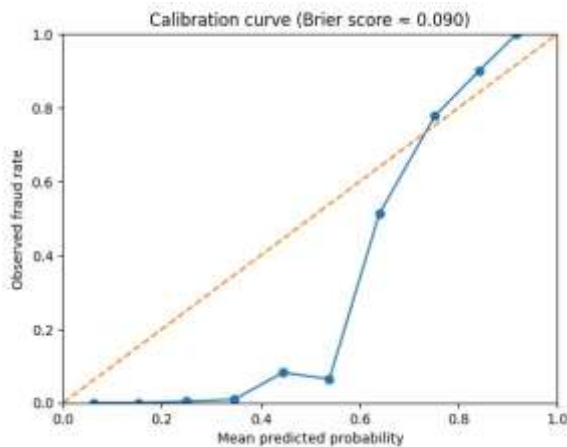
Fig 3: Bimodal Data Engineering for Insurance Fraud Detection: Integrating Fractal-Encoded Numerical Processing and NLP-Based Linguistic Structuring

6. Model Development and Evaluation

The model development for attack detection is conducted in three high-level stages. The first stage involves the exploration of various deep learning architectures and the selection of promising candidate models suitable for classifying fraud. The second stage covers the definition of training protocols for the identified candidates, an assessment of potential overfitting issues, and the identification of appropriate evaluation metrics as well as a validation plan. The third stage includes a concrete plan for the deployment of the best-performing model into continuous operation.

While deep learning discriminative classifiers are the core candidate architectures, some foundational supervised models have been added and will be evaluated to better understand the data and to provide performance benchmarks. In addition to convolutional networks for images, long short-term memory models for sequences, and spatiotemporal architectures (convolutional long short-term memory networks), it is also planned to explore less common unsupervised / self-supervised / adversarial approaches: generalised autoencoders, GANs and VAEs for both image and time series data anomalies. Predictive deep-fake models and GANs in particular might also indicate impression-manipulating sophistication or socially-aware generation capability. Other less-well-explored strategies like knowledge-distilled models – operationally efficient classifiers derived from knowledge-embodied and heavy transformers – will also be trialled on dataset subsets. Evaluation metrics will include standard classification measures (precision, recall, F1-score, ROC-AUC, Brier score) but with added sensitivity to the criticality of false negatives. For instance, a relatively high false-positive-to-false-negative cost ratio could justify maximising the false-negative-and-missed-class rate instead of minimising them. Models will also be cross-validated using multiple

techniques: standard k-fold, hold-out retroactive sets and time-respecting splits. Attention-based architectures will undergo attention-visualisation techniques to uncover which parts of the data are being used for classification. Sensitivity and explainability of supervised algorithms will be studied to understand what features they are using to identify fraud rapidly and which are “distracting” them during inference.

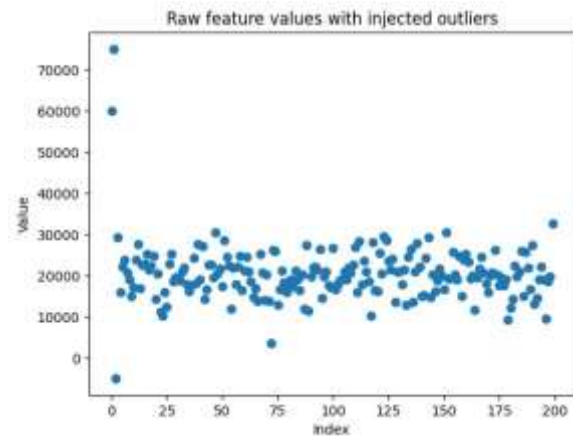


6.1. Model architectures and algorithms

A variety of model architectures is available for building classification models in insurance fraud detection. Popular supervised learning methods include tree-based algorithms such as random forests, gradient boosting, and XGBoost. Novel unsupervised techniques for fraud detection use autoencoders, deep support vector machines, and graph embedding methods. Unsupervised techniques facilitate novel insights while avoiding known pitfalls of supervised learning, such as the bias introduced by labelling and the dependence on correctly capturing fraud signals in the training data. Moreover, fraud, in particular new types, tends to occur very rarely in insurance data, making supervised learning less suited to circumstances where only historical data can be used.

An alternative approach to avoiding labelled training data is to derive inherent costs from the data and build cost-sensitive classifiers, capable of operating in an environment in which fraud occurs. For controlling misconduct, a new

way to optimize rules by minimizing label-dependent measures with a fast approximation using a reproducible cost-sensitive forest model has emerged. The new rule-induced cost-sensitive model provides an approximate solution to cost-sensitive multi-class problems while combining the advantages of tree-based methods and rule models. Finally, multiple-instance learning approaches tackle the issues stemming from the absence of precise labels and use the complement to classification for building explainable models in complex situations. All methods are equally suited for risk communication and should thus be applied in conjunction with interpretable AI methods.



Equation 3: Bitwise NOT label transform (as described)

“fraud value was inversely coded via bitwise not and then multiplied by -1 ”

For an integer y , bitwise NOT has the identity:

$$\sim y = -y - 1$$

Then multiplying by -1 :

$$-(\sim y) = -(-y - 1) = y + 1$$

So the transformation “bitwise not then multiplied by -1 ” is:

$$y' = -(\sim y) = y + 1$$



6.2. Evaluation metrics and validation

Different metrics communicate different information about the response of models to classified data, and thus, multiple metrics should be defined—relying on just one of them might yield a deceptive response. The metrics used in the development of the models include precision, recall, F1 score, area under receiver operating characteristic and calibration curves, and others. These are sufficient for evaluating whether the models are any good, as well as for suitable comparative analysis.

Precision indicates the percentage of true positives in relation to the number of positives classified by the model. Recall reveals the percentage of true positives in relation to the number of actual positives. The F1 score is the harmonic mean of precision and recall and combines precision and recall into one single score. The area under the receiver operating characteristic curve (ROC-AUC) indicates how well the model differentiates fraudulent transactions from legitimate transactions—an AUC of 0.5 indicates poor performance, while an AUC of 1.0 indicates perfect performance. Calibration indicates how well the model reflects prediction probabilities. Calibration is more important when making risky decisions. The point where calibration curve crosses $y = x$ is known as the ideal state. A model with a good calibration curve is viewed as a trustworthy model.

Raw Value (₹)	Capped ($\mu \pm 3\sigma$)	Min–Max Scaled
60000	42608	1.00
75000	42608	1.00
-5000	-626	0.00
22943	22943	0.54
19327	19327	0.47
21477	21477	0.51
18263	18263	0.44
15730	15730	0.39
25700	25700	0.60
20111	20111	0.49

Raw Value (₹)	Capped ($\mu \pm 3\sigma$)	Min–Max Scaled
23589	23589	0.56
16894	16894	0.42

7. Conclusion

The presented work illustrated the theoretical underpinnings and a viable instance of DevOps-supported data engineering for cloud-native agentic AI systems, applied to insurance fraud detection. The architecture, automations, data preparation workflows, and modelling process were documented as a set of design patterns that enable scalable, reproducible, and well-governed data engineering around ML models. Future research will address the development of production-ready models and consider the implications of ML-supported fraud detection in the domain of insurance.

As for direction, automating data engineering for DevOps principles and practices facilitates large-scale adoption of agentic AI, but brings new issues to the fore. Technical automation must be paralleled by normative developments in data governance and monitoring of ML models. Security is a pillar of DevOps, dovetailing with existing regulatory efforts; data privacy emerges as a key consideration; and ethical ML calls for careful attention within organisations before deploying decision-making or risk-assessing models. The ultimate objective is to offer solutions for shoring up insurance against fraud through solutions that align the prevention of malfeasance with the prevention of unfair corruption of the insurance safety net intended for community support.

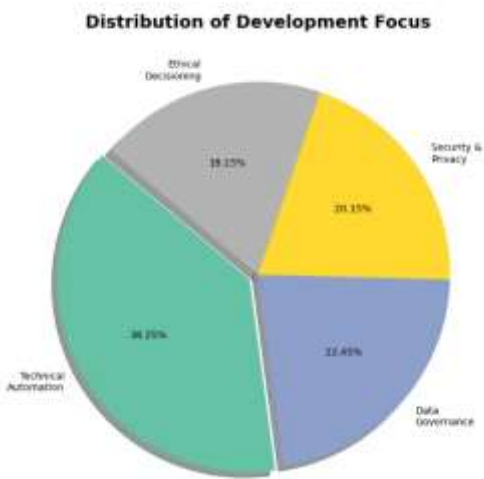


Fig 4: Distribution of Development Focus

7.1. Final Thoughts and Future Directions

Agentic AI impacts society and economics as intelligent, self-learning systems become ubiquitous. Insurance fraud detection is a prime candidate for expansive deployment because of the amount of data processed, the importance of avoiding financial losses through the detection of fraudulent claims, and the motivation of fraudsters to exploit system weaknesses.

The solution, based on a cloud-native DevOps approach to data engineering, integrates the core components of agentic AI in the continuous integration/continuous deployment of machine learning models. Security, compliance, data quality, and governance are continuously assured throughout the data lifecycle. These principles facilitate the scalable deployment of models tailored for the constant battle against ever-evolving fraudulent techniques, but scalability and operationalization remain to be proven. The commercial readiness of the system with DevOps-driven data engineering for agentic AI is underscored by the data preparation and ML pipelines demonstrating how cloud-native tools can continuously produce clean, well-structured data ready for feeding into ML algorithms and support the continuous training of a variety of models using both supervised and unsupervised techniques.

Continuous deployment of the ML component of agentic AI solutions in a dedicated ML area of the cloud, equipped with a model registry and constant data profiling, would enable the development of a diverse family of models focused on keeping pace with the fast-changing behavior of fraudsters without being caught in the risk–reward trap of real-world operationalization. For the industry, this solution represents not only a way to reduce costs and provide a better service to honest customers but also an opportunity for a real-time ML-driven solution for fraud detection, a follow-on level of thought that could enable customizability for other use cases in the insurance sector itself and the banking industry.

8. References

- Gomes, C., Jin, Z., & Yang, H. (2021). Insurance fraud detection with unsupervised deep learning. *Journal of Risk and Insurance*, 88(3), 591–624.
- Al-Hashedi, K. G., & Magalingam, P. (2021). Financial fraud detection applying data mining techniques: A comprehensive review from 2009 to 2019. *Computer Science Review*, 40, 100402.
- Jogi, N. L., Kumar, K. R., Mangala, N., M.Govindarajan, Gamini, P., & Bundhe, S. (2026). Machine Learning-Based Random Forest Model for Sustainable Supply Chain Sales Forecasting. In 2026 6th International Conference on Intelligent Technologies (CONIT) (pp. 1–6). IEEE. 2026 6th International Conference on Intelligent Technologies (CONIT).
<https://doi.org/10.1109/conit69683.2026.11620788>
- Alghofaili, Y., & Lukman, S. (2020). A financial fraud detection model based on LSTM deep learning technique. *Journal of Applied Security Research*, 15(4), 498–516.
- Aslam, F., Hunjra, A. I., Ftiti, Z., Louhichi, W., & Shams, T. (2022). Insurance fraud detection: Evidence from artificial intelligence and machine learning. *Research in International Business and Finance*, 62, 101744.

6. Recharla, M., Garapati, R. S., Bandi, V. D. V. K., Yandamuri, U. S., & Mangalampalli, B. M. (2026, March). Generative AI-Enhanced Data Engineering Pipelines for Predictive Biomarker Discovery in Alzheimer's and Kidney Disease. In 2026 IEEE International Conference on AI Engineering and Innovations (AIEI) (pp. 1-5). IEEE.
7. Bhowmik, A., Sannigrahi, M., Chowdhury, D., Dwivedi, A. D., & Mukkamala, R. R. (2022). DBNex: Deep belief network and explainable AI based financial fraud detection. In Proceedings of the 2022 IEEE International Conference on Big Data (pp. 3033–3042). IEEE.
8. Urunkar, A., Khot, A., Bhat, R., & Mudegol, N. L. (2022). Fraud detection and analysis for insurance claim using machine learning. In Proceedings of the 2022 IEEE International Conference on Signal Processing, Informatics, Communication and Energy Systems. IEEE.
9. Debener, J., Heinke, V., & Kriebel, J. (2023). Detecting insurance fraud using supervised and unsupervised machine learning. *Journal of Risk and Insurance*, 90(3), 743–768.
10. Lakhanpal, S., Tunjungsari, H. K., Raj, S., Chukka, A., Dawadi, D., & Mangalampalli, B. M. (2026, April). The Digital Transformation of Healthcare: Balancing Opportunities, Risks, and Ethical Considerations. In 2026 International Conference on Multidisciplinary Innovations For Smart & Sustainable Future (MISSF) (pp. 1-6). IEEE.
11. Chousa Santos, M., Pereira, T., Mendes, I., & Amaral, A. (2024). Machine learning for insurance fraud detection. In *Internet of Everything* (pp. 49–60). Springer.
12. Banulescu-Radu, D., Yankol-Schalck, M., & Kriebel, J. (2024). Practical guideline to efficiently detect insurance fraud in the era of machine learning: A household insurance case. *Journal of Risk and Insurance*, 91(4), 867–913.
13. Vorobyev, I. (2024). Fraud risk assessment in car insurance using claims graph features in machine learning. *Expert Systems with Applications*, 251, 124109.
14. Nandan, B. P., Kumar, M. V. K., Garapati, R. S., Bandi, V. D. V. K., Davuluri, P. N., & Mangalampalli, B. M. (2026, March). AI-Enhanced Semiconductor Yield Optimization Using Hybrid Deep Learning and Edge Data Analytics. In 2026 IEEE International Conference on AI Engineering and Innovations (AIEI) (pp. 1-6). IEEE.
15. Ding, N., Ruan, X., Wang, H., & Liu, Y. (2025). Automobile insurance fraud detection based on PSO-XGBoost model and interpretable machine learning method. *Insurance: Mathematics and Economics*, 120, 51–60.
16. Wang, Z., Chen, X., Wu, Y., Jiang, L., Lin, S., & Qiu, G. (2025). A robust and interpretable ensemble machine learning model for predicting healthcare insurance fraud. *Scientific Reports*, 15, 218.
17. Özaltın, Ö., & Karadağ Erdemir, Ö. (2025). Detecting automobile insurance fraud using a novel penalty-driven feature selection method with particle swarm optimization and machine learning classifiers. *Scientific Reports*, 15, 41793.
18. Gopinath, A., Davuluri, P. N., Kumar, U. B., MP, S., & Yamini, G. (2026, April). Communication Aware Malware Detection Using Network Flow Bytecode Fusion With Adaptive Focal Optimization. In 2026 2nd International Conference on Intelligent Systems and Computational Networks (ICISCN) (pp. 1-8). IEEE.
19. Kumar, P. A., & Sountharajan, S. (2025). Insurance claims estimation and fraud detection with optimized deep learning techniques. *Scientific Reports*, 15, 27296.
20. Khalil, A. A. (2026). Enhancing insurance fraud detection accuracy with integrated machine learning and statistical methods. *Computational Economics*, 68, 707–738.
21. Yankol-Schalck, M. (2026). Auto insurance fraud detection: Machine learning and deep learning applications. *Journal of Risk and Insurance*, 93, 534–568.

22. Davuluri, P. S. L. (2023). AI-Augmented Sanctions Screening: Enhancing Accuracy and Latency in Real Time Compliance Systems. *AI-Augmented Sanctions Screening: Enhancing Accuracy and Latency in Real Time Compliance Systems* (December 15, 2023).
23. Motie, S., & Raahemi, B. (2024). Financial fraud detection using graph neural networks: A systematic review. *Expert Systems with Applications*, 240, 122156.
24. Zhang, R., Cheng, D., Yang, J., Ouyang, Y., Wu, X., Zheng, Y., & Jiang, C. (2024). Pre-trained online contrastive learning for insurance fraud detection. *Proceedings of the AAAI Conference on Artificial Intelligence*, 38(20).
25. Bao, Y., Ke, Y., Li, Y., Tu, J., & Chen, W. (2020). Detecting accounting fraud in the big data era: A new framework for machine learning. *International Journal of Accounting Information Systems*, 38, 100455.
26. West, J., & Bhattacharya, M. (2020). Intelligent financial fraud detection practices: An investigation. *International Journal of Information Management*, 50, 170–186.
27. Kolla, S. K., Bandi, V. D. V. K., & Meda, R. (2026). Comment on “From laboratory promise to decision-grade practice: strengthening reproducibility and casework relevance in AI-assisted bloodstain ageing”. *Forensic Science, Medicine and Pathology*, 1-2.
28. Hilal, W., Gadsden, S. A., & Yawney, J. (2022). Financial fraud: A review of anomaly detection techniques and recent advances. *Expert Systems with Applications*, 193, 116429.
29. Ali, A., Pfohl, S., & Bader, M. (2022). Financial fraud detection based on machine learning and deep learning techniques: A systematic review. *Applied Sciences*, 12(19), 9637.
30. Fiore, U., De Santis, A., Perla, F., Zanetti, P., & Palmieri, F. (2021). Using generative adversarial networks for improving classification effectiveness in credit card fraud detection. *Information Sciences*, 479, 448–455.
31. Jurgovsky, J., Granitzer, M., Ziegler, K., Calabretto, S., Portier, P. E., He-Guelton, L., & Caelen, O. (2020). Sequence classification for credit-card fraud detection. *Expert Systems with Applications*, 100, 234–245.
32. Lucas, Y., Portier, P.-E., Laporte, L., Caelen, O., He-Guelton, L., Granitzer, M., & Calabretto, S. (2020). Multiple perspectives HMM-based feature engineering for credit card fraud detection. *Pattern Recognition*, 104, 107330.
33. Carcillo, F., Le Borgne, Y.-A., Caelen, O., & Bontempi, G. (2021). Combining unsupervised and supervised learning in credit card fraud detection. *Information Sciences*, 557, 317–331.
34. Carcillo, F., Le Borgne, Y.-A., Caelen, O., Kessaci, Y., Oblé, F., & Bontempi, G. (2021). Combining unsupervised and supervised learning in credit card fraud detection. *Information Sciences*, 557, 317–331.
35. Randhawa, K., Loo, C. K., Seera, M., Lim, C. P., & Nandi, A. K. (2020). Credit card fraud detection using AdaBoost and majority voting. *IEEE Access*, 6, 14277–14284.
36. Kolla, S. H., Nagabhyru, K. C., & Kummari, D. N. (2026). Letter to the Editor re: “CurvAssist: An AI assisted pipeline for penile shaft segmentation/curvature measurement in children with hypospadias”. *Journal of Pediatric Urology*.
37. Xuan, S., Liu, G., Li, Z., Zheng, L., Wang, S., & Jiang, C. (2020). Random forest for credit card fraud detection. In *Proceedings of the IEEE International Conference on Machine Learning and Applications*. IEEE.
38. Bhatla, T. P., Prabhu, V., & Dua, A. (2020). Understanding credit card fraud detection using machine learning techniques. *International Journal of Computer Applications*, 175(8), 15–20.
39. Weber, M., Domeniconi, G., Chen, J., Weidele, D. K. I., Bellei, C., Robinson, T., & Leiserson, C. E. (2020). Anti-money laundering in Bitcoin: Experimenting with graph convolutional networks for financial forensics. *Proceedings of the 29th ACM International Conference on Information and Knowledge Management*, 1461–1470.
40. Wang, J., Wen, R., Wu, C., Huang, Y., & Xiong, J. (2021). FdGars: Fraud detection in graph-based

- account-relation systems. Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, 411–419.
41. Dou, Y., Liu, Z., Sun, L., Deng, Y., Peng, H., & Yu, P. S. (2020). Enhancing graph neural network-based fraud detectors against camouflaged fraudsters. Proceedings of the 29th ACM International Conference on Information and Knowledge Management, 315–324.
 42. Subramanian, N., & Kolla, T. Equity in Healthcare Access Policy Lessons from Universal Coverage Models.
 43. Liu, Z., Dou, Y., Yu, P. S., Deng, Y., & Peng, H. (2021). Alleviating the inconsistency problem of applying graph neural networks to fraud detection. Proceedings of the 30th ACM International Conference on Information and Knowledge Management, 1563–1572.
 44. Wang, D., Lin, Z., & Cui, P. (2021). A survey on graph neural networks for fraud detection. ACM Computing Surveys, 54(4), 1–36.
 45. Kipf, T. N., & Welling, M. (2020). Semi-supervised classification with graph convolutional networks. International Conference on Learning Representations.
 46. Lundberg, S. M., & Lee, S.-I. (2020). A unified approach to interpreting model predictions. Advances in Neural Information Processing Systems, 30.
 47. Molnar, C. (2022). Interpretable machine learning: A guide for making black box models explainable. Independently published.
 48. Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Benetot, A., Tabik, S., Barbado, A., García, S., Gil-López, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. Information Fusion, 58, 82–115.
 49. Meske, C., Bunde, E., Schneider, J., & Gersch, M. (2022). Explainable artificial intelligence: Objectives, stakeholders, and future research opportunities. Information Systems Management, 39(1), 53–63.
 50. Bauer, K., Weitz, K., & Schütte, R. (2023). Explainable artificial intelligence for financial fraud detection: A user-centered approach. Finance Research Letters, 58, 104309.
 51. Vijayanand, D., & Smrithy, G. S. (2025). Explainable AI-enhanced ensemble learning for financial fraud detection in mobile money transactions. Journal of Information Science and Engineering, 19(1).
 52. Schick, T., Dwivedi-Yu, J., Dessi, R., Raileanu, R., Lomeli, M., Hambro, E., Zettlemoyer, L., Cancedda, N., & Scialom, T. (2023). Toolformer: Language models can teach themselves to use tools. Advances in Neural Information Processing Systems, 36.
 53. Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing reasoning and acting in language models. International Conference on Learning Representations.
 54. Wang, L., Ma, C., Feng, X., Zhang, Z., Yang, H., Zhang, J., Chen, Z., Tang, J., Chen, X., Lin, Y., Zhao, W. X., Wei, Z., & Wen, J.-R. (2024). A survey on large language model based autonomous agents. Frontiers of Computer Science, 18, 186345.
 55. Wu, Q., Bansal, G., Zhang, J., Wu, Y., Li, B., Zhu, E., Jiang, L., Zhang, X., Zhang, S., Liu, J., Awadallah, A. H., White, R. W., Burger, D., & Wang, C. (2024). AutoGen: Enabling next-gen LLM applications via multi-agent conversation. Proceedings of the 2024 International Conference on Learning Representations Workshops.
 56. Liu, X., Yu, H., Zhang, H., Xu, Y., Lei, X., Lai, H., Gu, Y., Ding, H., Men, K., Yang, K., Zhang, S., Deng, X., Zeng, A., Du, Z., Zhang, C., Shen, S., Zhang, T., Su, Y., Sun, H., Huang, M., Dong, Y., & Tang, J. (2024). AgentBench: Evaluating LLMs as agents. International Conference on Learning Representations.
 57. Zhou, Y., Muresanu, A. I., Han, Z., Paster, K., Pitis, S., Chan, H., & Ba, J. (2023). Large language models are human-level prompt engineers. International Conference on Learning Representations.
 58. Yao, S., Yu, D., Zhao, J., Shafran, I., Griffiths, T. L., Cao, Y., & Narasimhan, K. (2023). Tree of thoughts: Deliberate problem solving with large language



- models. *Advances in Neural Information Processing Systems*, 36.
59. Mialon, G., Dessì, R., Lomeli, M., Nalmpantis, C., Pasunuru, R., Raileanu, R., Rozière, B., Schick, T., Dwivedi-Yu, J., Celikyilmaz, A., Grave, E., & Scialom, T. (2023). Augmented language models: A survey. *Transactions of the Association for Computational Linguistics*, 11, 119–142.
 60. Zhou, Y., Shen, Y., Li, H., & others. (2024). Large language models for autonomous agents: Architectures, applications, and challenges. *ACM Computing Surveys*, 56(8).
 61. Zaharia, M., Chen, A., Davidson, A., Ghodsi, A., Hong, A., Konwinski, A., Murching, S., Nykodym, T., Ogihara, M., Parkhe, M., Xie, F., & Zumar, C. (2023). The data-centric AI manifesto. *Communications of the ACM*, 66(8), 36–43.
 62. Kreuzberger, D., Kühn, N., & Hirschl, S. (2023). Machine learning operations (MLOps): Overview, definition, and architecture. *IEEE Access*, 11, 31866–31879.
 63. Polyzotis, N., Roy, S., Whang, S. E., & Zinkevich, M. (2021). Data lifecycle challenges in production machine learning: A survey. *ACM SIGMOD Record*, 49(2), 17–28.
 64. Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., Chaudhary, V., Young, M., Crespo, J.-F., & Dennison, D. (2020). Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems*, 28.