



Adaptive Agentic AI for Cloud-Native Fraud Detection in InsurTech

James Robertson
Asst Professor

Abstract

The pandemic-induced surge in encompass an explosion of opportunities for fraud and thus the need for improved fraud-detection analytics. When external stakeholders such as financial service providers and law enforcement agencies are involved, fraud detection and prevention systems must go beyond detection alone. Agentic AI enables an improved intelligent-agent-centered design for InsurTech and provides a framework for decisions involving an increase in risk. The hybridization of agentic AI with cloud-native concepts further extends the design pattern to fraud-detection platforms developed on cloud services. Emerging specifications of the fraud detection domain call for agentic systems that take such factors into consideration. While theory enables advances in design and artifact construction, novel system designs clarify operation and integration within the broader operation of an InsurTech system. Together they offer a foundation for future real-world implementations of agentic concepts that can deliver business value.

Research avenues covering core themes facing agents and cloud-based design pattern components guide practical implementation and further development of the proposed concepts, including pioneering cloud-based fraud detection for InsurTech executives. Reactions, responsibility, evolution, and privacy concern the decision space of agent-based privacy-preserving systems, and linked supporting factor rails define considerations inherent in managing the design and construction of high-quality aged liquor. Patterns facilitate cloud-based fraud-detection platforms capable of managing application, data, and model service risk transparently and efficiently throughout their operating life cycle. The specification of continuous integration and delivery pipelines serves as a guide for constructing fraud models and related assets.

Keywords: Agentic Artificial Intelligence, InsurTech Fraud Detection, Cloud-Native Fraud Platforms, Intelligent Agent Design, Hybrid AI Architectures, Pandemic-Driven Fraud Risk, Privacy-Preserving AI Systems, Risk-Aware Decision Frameworks, Financial Crime Prevention, Autonomous Fraud Mitigation, Multi-Stakeholder Fraud Ecosystems, Cloud-Based Analytics, Continuous Integration And Delivery Pipelines, Model And Data Governance, Responsible AI Systems, Scalable Fraud Analytics, Real-Time Fraud Prevention, Secure Financial Platforms, AI-Driven Risk Escalation, Lifecycle-Aware Fraud Management.

1. Introduction

Many InsurTech platforms are falling victim to fraud through phishing, spam, and account takeover attacks. Fraud often manifests in user behavior that undermines fraud detection, risk prevention, and operational optimization. Insurers must spend an increasing share of their earned premiums on fraud remediation—a volume that can exceed 100% of the firm's operating profits. Ongoing economic turmoil and the pandemic have spurred

insurers to automate and streamline business processes, but investment priorities are shifting, with Crime and Fraud Risk Management budgets currently in decline. Criminals are well aware of the cross-industry migration to online services, and InsurTech still has some of the lowest budgets allocated to cybercrime countermeasures. Malicious insiders, account takeovers, and collusion account for a disproportionately large share of online fraud. Standard anti-fraud measures work but are costly and may create stickiness in the product that is harmful to the



business. In response, an agentic AI encompasses agent Autonomy, human Interaction, Trust, and providing capabilities to derive Governance and Compliance. The fraud landscape in InsurTech deserves special attention, and a broader perspective proposes requirements for a cloud-native platform including Continuous Integration and Delivery for Machine Learning Models, Data Virtualization, and Federated Learning.

1.1. Overview of the Study

Fraud in InsurTech is receiving increasing attention due to fast-growing premium volumes and expansion of online channels, with regulators deploying specific reporting requirements. However, most modern InsurTech platforms do not monitor transaction integrity, with existing probative models unsuited for production. Although periodic auditing serves compliance, a responsive detection solution that limits operational and reputational risks remains absent. This study recognizes such a gap and advances both the technology and its methodological evaluation. The evaluation employs a fraud dataset created from a major study in cyber deception. Fraud agents dynamically deploy various deception tactics across different products over time, eliciting a complete spectrum of fraud scenarios. Agentic decision-making serves as a proxy for missing detection models. The provided solution is adaptable, accounts for temporal changes, and employs risk scoring backed by robust model management pipelines. The restructuring supports federation across InsurTech platforms and partners while addressing privacy concerns.



Fig 1: Adaptive Fraud Governance in InsurTech: A Federated Risk-Scoring Framework for Real-Time Detection of Dynamic Cyber Deception and Agentic Fraud Tactics

2. Background and Motivation

Fraud is a pressing issue in InsurTech, where it has reached industry-critical levels. Recent developments in market regulation both drive the need for action and increase the complexity of detection efforts. Agentic AI offers a novel solution for adaptive detection, yet deployment remains pending; cloud-native architecture provides a basis for support. Recent months have witnessed the continued surge of agentic AI systems into enterprise environments, driven by open-source implementations and cloud deployment services. The dangers of poor implementation are also becoming evident, underscoring the need for fundamental theory and model specification for the domain. Agentic systems are different from chatbots, which are just tools without the ability to make decisions on their own. Agentic capabilities allow systems to make decisions with the expertise of a risk manager instead of being explicitly programmed into a decision tree. This kind of AI is especially useful for fraud detection, as it requires

fraud mitigation, detection, and real-time management while maintaining the required governance and compliance controls.

The focus of research and design in fraud detection has primarily been around the construction of machine learning models to identify aberrant behaviours and scoring of transactions that are not explicitly supervised. Typically, these models and scoring methodologies are not considered in conjunction with traditional risk-scoring methodologies, nor have they been expressly defined for real-time decisions during transaction flows. The provision of these services is neither autonomous nor controlled by a specific business function. Consequently, they do not integrate with existing governance and compliance frameworks. An agentic approach provides the foundations to fill these gaps while enabling deeper examination of the underlying technologies.

3. Theoretical Foundations of Agentic AI in Fraud Detection

The core theories, models, and ontologies supporting agentic AI in fraud-detection tasks are presented. Agent capabilities, negotiation with humans, accountability, explainability, and interoperability frameworks are defined. The Utility Theory (Von Neumann & Morgenstern), Model of E-Government (Oberholzer et al.), Agency Theory (Eisenhardt, 1989), Agent-Oriented Systems (Varun & Singh, 2018), E-Rulemaking (Oberholzer et al.), Explanation in AI and ML (Gilpin et al.; Ribeiro et al.), and the Ontology for Agent Capability Specification (US Dept of Defense) lay the groundwork for all types of agent-based detection systems. Agentic AI departs from traditional agent-oriented perspectives by incorporating an explicit definition of the type of intelligence and autonomy offered to agents, along with the associated types of decisions. Detection techniques that are powered by agentic AI commonly have autonomous capabilities restricted at certain stages of the detection process. Figure 2 provides a detailed explanation of the dimensions defining those capabilities. Although certain decisions must ultimately be

taken by a human agent, these techniques allow the lower-level agents to operate autonomously, interacting with their human counterparts when needed for negotiation, support, or specification of particular roles. In this context, it is also important to have a framework that ensures accountability of AI-driven systems.

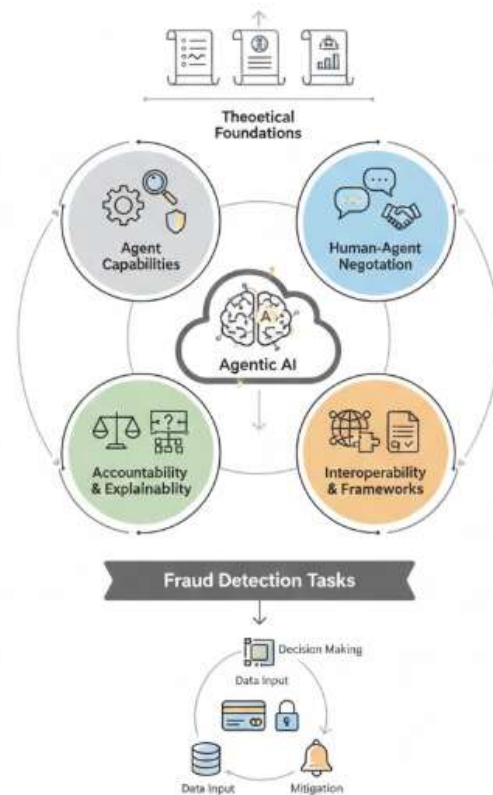


Fig 2: Governing Autonomy: A Multi-Theoretical Framework for Agentic AI in Fraud Detection, Accountability, and Human-Agent Negotiation

3.1. Key Concepts and Frameworks for Agentic AI in Fraud Detection

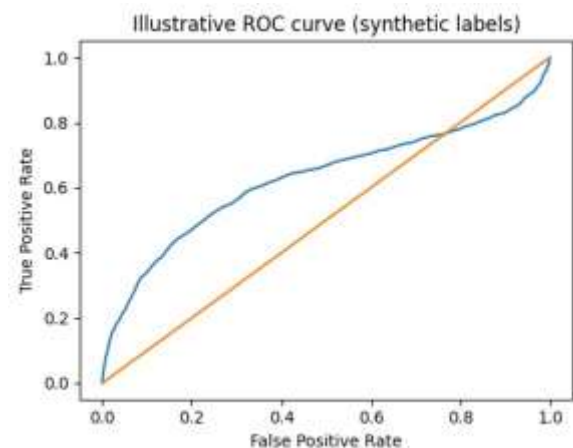
Four concepts are instrumental for embedding an agent in fraud detection: the capabilities that enable its interactions with humans; the ability to negotiate with users, particularly for critical actions; a framework for

governance and accountability; and a mechanism that supports the explanations required in regulated environments with human oversight.

Autonomous agent capabilities must support the execution of both reactive and pro-active fraud-detection actions; that is, a detection alert or the execution of a pre-trained detection policy can be performed automatically, while the execution of an untrained detection policy for individual at-risk users must involve user feedback. In particular, the analysis of the importance of predictions for well understood models and the identification of suspected fraud have an explicit user feedback requirement. For clear and interpretable models, this feedback can be provided by crowd sourcing.

Even well governed agents with established and robust decision boundaries require bouts of user negotiation where fraud detection is concerned, particularly in regulated environments. This negotiation need not be recognized and incorporated in the operational design framework of the agent. However, during its long term development and evaluation phases, the agent must be seen not merely as a reactive entity but also as one capable of initiating negotiation with humans.

To conform with the governance demands of users and regulators, and to enable simple human-agent communication, a framework within which accountability is clearly established, and its provision minimised, must be specified. The accountability framework defines when an agent must be accountable for its actions, and to what extent it successfully provides such accountability. The latter depends on the sufficiency of the principles and justifications offered both by the agent and the learning algorithms it employs.



4. Architectural Principles for Cloud-Native Deployments

In cloud-native environments, three architectural objectives are paramount: scale—particularly in terms of traffic fluctuations—resilience to component or service outages, and security. Achieving these goals requires attention throughout the design, but especially during decomposition into modular components. A key design capacity introduced by the cloud paradigm is the ability to decompose tasks into small units of work that can be executed rapidly and efficiently in geographically-distributed locations. These smaller tasks can represent microservices, each capable of executing an atomic capability independently in response to external requests, encapsulating the required functions within service boundaries. Cloud providers also enable the serverless paradigm, allowing developers to define functions that are executed in response to demand, either directly in response to event notifications or triggered by monitoring conditions in other services. Wrapping reusable functionality—such as model scoring—as a serverless function combined with a supporting data preparation phase often costs significantly less than hosting the same model within the context of a full-fledged microservice.

Pattern-based design enables a more general reuse of class libraries, utility code, and infrastructure templates when



implementing the microservices and supporting architectures. Some patterns—particularly cloud watch tower, cosmic petri nets, terraform template factories, and flickering static lamps—help to turn architectural goals into concrete scalable infrastructure, others—such as component watchful eyes, wicked telemetry plots, and continuous portraiture—provide insights into how the operational environment is running and when human intervention might be required. Others exploit these tracking patterns to enable a “shift left” philosophy to monitoring and management: providing early warning of nearing capacity limits, failures in prior executions—such as model scoring—insufficient quality of operating fragments, and operational threats.

Equation 2: From risk score to decisions: thresholds + expected utility (agentic action)

2.1 Action set

Let the agent choose an action $a \in \mathcal{A}$, e.g.

- a_0 : allow
- a_1 : step-up authentication
- a_2 : hold + manual review
- a_3 : block/blacklist (high-impact)

2.2 Expected utility (step-by-step)

Let $y \in \{0,1\}$ be the true class (0 legit, 1 fraud).
Let $\hat{p}_t = \Pr(y = 1 \mid \mathbf{x}_t)$. In our system, $\hat{p}_t$ is the calibrated risk score r_t .

Define utilities (or costs):

- $U(a, y)$: utility if we take action a and truth is y

Expected utility of action a :

$$\mathbb{E}[U(a, y) \mid \mathbf{x}_t] = U(a, 1)\hat{p}_t + U(a, 0)(1 - \hat{p}_t)$$

Optimal action:

$$a_t^* = \operatorname{argmax}_{a \in \mathcal{A}} (U(a, 1)\hat{p}_t + U(a, 0)(1 - \hat{p}_t))$$

2.3 Threshold form (common special case)

If you only have two actions: allow vs block, with costs:

- false negative cost C_{FN} (fraud allowed)
- false positive cost C_{FP} (legit blocked)

Choose “block” when:

$$C_{FP}(1 - \hat{p}_t) < C_{FN}\hat{p}_t$$

Solve step-by-step:

$$C_{FP} - C_{FP}\hat{p}_t < C_{FN}\hat{p}_t \quad C_{FP} < (C_{FN} + C_{FP})\hat{p}_t \quad \hat{p}_t > \frac{C_{FP}}{C_{FN} + C_{FP}} = \tau$$

So decision rule:

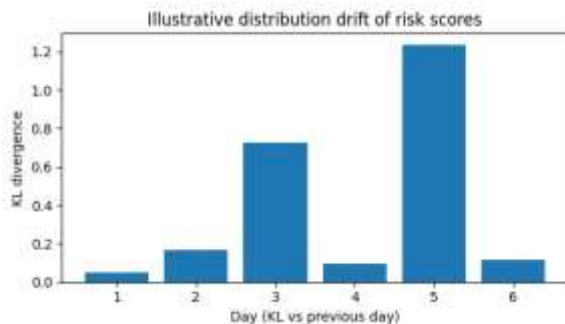
$$\text{block if } r_t > \tau$$

4.1. Microservices and Serverless Constructs

Practical cloud-native deployments decompose detection functionality into smaller components, capitalize on cloud facilities for scaling and resilience, and make event-driven interaction a natural fit. Control- and management-oriented elements, however, benefit from analytical safety and governance strategies, and strong confidentiality requirements of fraud detection often favour on-premise operations. Consequently, modules fall into three categories. Microservices that can scale independently handle core detection functionality (e.g., real-time risk scoring, anomaly detection, data drift checks, behavioral profiling). These interact asynchronously via message queues that enable decoupling and resilience. Non-analytical components—risk-model scoring, risk-color-articulating components, historical false-positive analyses, and machine-learning models—are reduced to a virtual-service lookup with all heavyweight operations abstracted away as black boxes. The model-development pipelines for training and validating these are outlined next; aspects can later be refined for additional control and governance. Any



remaining cloud-based elements are served by serverless constructs rather than more heavyweight microservices, although the latter would remain viable. Event-driven microservices can be conveniently implemented as serverless components in environments offering the necessary facilities. The arguments for such an approach—increased speed of delivery and reduced maintenance burden—are therefore re-emphasized here. Serverless components are coded without explicit concern for infrastructure provisioning, maintenance, capacity planning, scaling, or high availability—all facilities being provided automatically (whether by the platform or by underlying cloud services) and charged on an as-used basis. The risk that a burst of concurrent executions might exceed a service provider’s limits and thus cause performance degradation is therefore major; as a result, applications that might hit that risk naturally decompose into separately offering validation cycles without strong tight-coupling characteristics.



5. Adaptive Fraud Detection: Methods and Models

Fraud detection is typically framed as a supervised classification task using historical transaction metadata, but such approaches can struggle in practice due to small positive samples. To alleviate these issues, detection is approached in a continuous risk-scoring framework that makes use of three complementary perspectives: scoring models that rank risk of individual transactions as they happen, an operational risk-monitoring setup to map the real-time situation of the transaction pool, and an anomaly-

detection approach complemented with user-behavioural profiling. Each perspective offers different insights about fraud-related risk and helps build stronger protection against loss.

As in any risk-based approach, the main objective is to identify and prioritize transactions that are potentially risky and deserve further inspection or ad hoc mitigating actions. The difference is that a number of assumptions that make traditional classification models work may no longer hold—namely that enough flagged cases exist to build a predictive model and that the classification decisions are fully automated. Ultimately, the techniques converge on the same notion—an augmenting-and-adapting framework that prioritizes risk in near real-time rather than trying to make instant Go/No-Go decisions based on potentially flawed classification models.

Stage	Key checks	Typical artifact
Ingest	Schema, freshness, access control	Raw parquet snapshot
Validate	Missingness, leakage, label quality	Data quality report
Feature build	Versioned features, reproducibility	Feature store commit
Train	Re-train with fixed seeds & configs	Model weights + card
Evaluate	Holdout + time-split CV + bias checks	Eval report + thresholds
Deploy	Canary/blue-green + policy gates	Serving endpoint
Monitor	Drift, latency, FP/FN, cost	Dashboards + alerts
Rollback	Auto revert on KPI regression	Pinned prior model

5.1. Real-Time Risk Scoring

Real-time risk scoring aims to perniciously label received events, with a particular focus on the threat of fraud. Though pure classification models for identifying fraud instances are commonplace, here the temporal



characteristics of fraud are key. Low-quality fraud, often in the form of small value or denied credit card payments following stolen card data releases, arises in real time and manifests in sudden bursts. Therefore, the objective is not simply to distinguish good from bad at a single moment but to classify risk over a finite time horizon. Fraudsters look for new targets and will go through several in one session; if they are successful, other fraudsters will come in seeing their success via the social channels. In contrast, an agent generally wants to check the same base over time, hence the appearance of a diverging Horman(VAR) in extreme cases. These characteristics yield the highest HY rates. Predictive features are therefore time-dependent and predictive models must follow suit. The underlying time series data are event counts, these series exhibit strong clustering, both temporally (time-in-bracket) and spatially (place-in-world). Temporal aspects are addressed through the natural log of the counts at site t converted to -0.5count , this gives a good ‘not seen yet’ heuristic over possible estimates. The Voges-Massel series with a symbolic scale is tested through explaining the behaviour of the Anglosperg. It acts as the ideal base model but being construction counts, which it tends to smooth and magnify. A custom KNN approach best reproduces the Banknock-normalised Braille series, and simple vector techniques applied to renewed counts produce large and interpretable word-above-empty scores.



Fig 3: Temporal Dynamics and Spatio-Temporal Clustering in Real-Time Risk Scoring: A Predictive Framework for High-Frequency Fraud Detection

5.2. Anomaly Detection and Behavioral Profiling

A combination of anomaly detection (AD) and behavioral profiling (BP) is well-poised to capture unexpected behavioral variation, whether within or among individuals. Anomaly detection identifies atypical activity across the whole group of individuals and is capable of surfacing fraud at any decision point. Behavioral profiling serves as a foundation for anomaly detection by tracking the status quo of user behavior, enabling the detection of deviations with insight about the change.

For detecting fraud at specific decision points through AD, traditional multivariate techniques (e.g., PCA, covariate, or distributional approaches) support the identification of atypical behavior. User groupings must then be carefully defined to capture genuine divergence; for example, geographically oriented groupings can be useful when assessing insurance claims that are highly sensitive to adverse weather events. Static profiling can detect sudden spikes by establishing intrusion detection thresholds; however, these simple rules are prone to false positives for “legitimate outliers” (e.g., people traveling). Hence, sufficient historical data need to be available for anomaly detection to be effective. Addressing such horizontal variation can be done by using filtering or grouping approaches.

Reaction Model Drift Handling is a key mechanism to operationalize drift concepts. A combination of online learning to automatically update the profile and a detector to indicate the need for retraining when the underlying distribution changes can address conventional baseline-drift challenges. To improve learning stability, the underlying distribution may be decoupled from the prediction by estimating and calibrating via a stratification on segments of interest. Such temporal stratification can also improve detection performance through deeper temporal significance, cross-validation regularization, and drift-targeted learning of non-event periods.

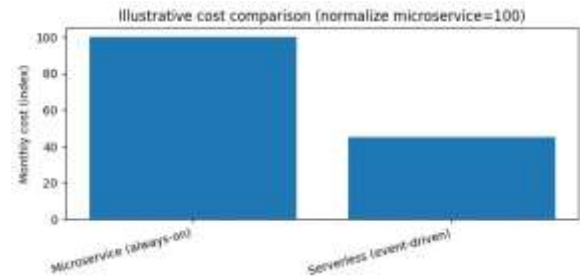


Metric	Example (baseline model)	Example (adaptive agentic model)
Precision	0.72	0.75
Recall	0.55	0.68
F1	0.62	0.71
False Positive Rate	0.030	0.028
Latency p95 (ms)	180	95

6. Agentic Capabilities in InsurTech

Agentic AI has significant potential to influence the mitigation of fraud in InsurTech by integrating into the workflows of insurance providers and stakeholders such as regulators or data partners. Insurance agents are responsible for all facets of the insurance lifecycle. Although much of their work is prescriptive, involving processes, forms, and rules, some of their tasks require making decisions about unusual situations that fall outside typical boundaries. Many of these decisions can potentially affect fraud detection,

Fraudulent claims typically fall outside the normal operation of an insurance provider, and human agents may decide to reject a claim as being suspicious or to notify the insurer's fraud team. These decisions differ from those about non-suspicious claims. Managing fraud detection requires external awareness, which presents another point at which human agents and detection systems must coordinate. Thus, the optimal level of decision autonomy and governance can differ between these points, and these aspects must be clearly specified. Incorporating these factors into the system can increase its acceptance and effectiveness by building—rather than undermining—human accountability.



Equation 3: Anomaly Detection + Behavioral Profiling (AD + BP)

3.1 Behavioral profile model

For each user u , define a profile distribution of behavior features (e.g., time-of-day, amount, device, location):

$$\mathbf{x}_{u,1}, \dots, \mathbf{x}_{u,n}$$

Estimate mean and covariance:

$$\begin{aligned} \boldsymbol{\mu}_u &= \frac{1}{n} \sum_{k=1}^n \mathbf{x}_{u,k}, \quad \boldsymbol{\Sigma}_u \\ &= \frac{1}{n-1} \sum_{k=1}^n (\mathbf{x}_{u,k} - \boldsymbol{\mu}_u)(\mathbf{x}_{u,k} - \boldsymbol{\mu}_u)^\top \end{aligned}$$

3.2 Deviation score (Mahalanobis distance)

For a new transaction $\mathbf{x}_{u,t}$:

$$D^2(\mathbf{x}_{u,t}) = (\mathbf{x}_{u,t} - \boldsymbol{\mu}_u)^\top \boldsymbol{\Sigma}_u^{-1} (\mathbf{x}_{u,t} - \boldsymbol{\mu}_u)$$

Flag anomaly if:

$$D^2(\mathbf{x}_{u,t}) > \chi_{d,1-\alpha}^2$$

(χ^2 threshold for dimension d , significance α).

3.3 Population-level anomaly detection via PCA residual

Let $\mathbf{x} \in \mathbb{R}^d$. Fit PCA on “normal” data, keep k components. Reconstruction:

$$\hat{\mathbf{x}} = \mathbf{P}_k \mathbf{P}_k^T \mathbf{x}$$

Residual:

$$e = \|\mathbf{x} - \hat{\mathbf{x}}\|_2$$

Flag if $e > \eta$ (threshold from validation).

6.1. Decision Autonomy Levels and Governance

Cloud-Native Agentic AI for Adaptive Fraud Detection in InsurTech Platforms

Agentic AI can operate within insurers' core workflows and comply with industry regulations while checking for and counteracting fraud. Decision autonomy levels, boundaries, human-in-the-loop mechanisms, and auditing and compliance paths indicate how far AI can act on its own in any fraud-detection task. Governance principles guide oversight and control.

Depending on context and risk appetite, an organization can decide how much fraud-related work to leave to AI. Within insurance, the cost and risk of human regularization policy violations makes auditing a special concern, calling for deeper investigation. The trade-off between operational accuracy and effective fraud management introduces further complexity. While detection and prevention tasks typically benefit from light human involvement, correction inevitably requires human interaction. AI can determine, when assigning risk scores in real time, whether to take action or escalate the case for human review.

7. Cloud-Native Platform Design Patterns

Reusable patterns enable scalable fraud-detection platforms in cloud environments. Fraud-detection models typically require regular retraining on fresh data, and the modelling pipeline must address data availability, quality, and relevance. Continuous Integration and Continuous Delivery (CI/CD) is thus an important foundational capability for such models and is an established practice across mainstream machine learning. However, the iterative pipelines require adaptation for effective application in fraud detection. Other patterns, such as data virtualisation

and federated learning, support cross-partner collaboration, enabling improved detection while ensuring privacy compliance according to domain regulations.

1. Continuous Integration and Delivery for Fraud Models: Fraud-detection models typically require reraining on fresh data on a non-trivial schedule and are thus enabled by Continuous Integration and Continuous Delivery (CI/CD) mechanisms for model training, validation, deployment, rollback and monitoring. CI/CD generally integrates the following tasks: pull new training data, apply MLops best practices to clean and preprocess, ensure a valid feature set, retrain the model, validate against a holdout set, deploy if the new model outperforms the old one, monitor performance, and trigger rollback if necessary. Notwithstanding these core tenets of CI/CD, the pathways require careful adaptation to be effective and efficient in a fraud-related context.

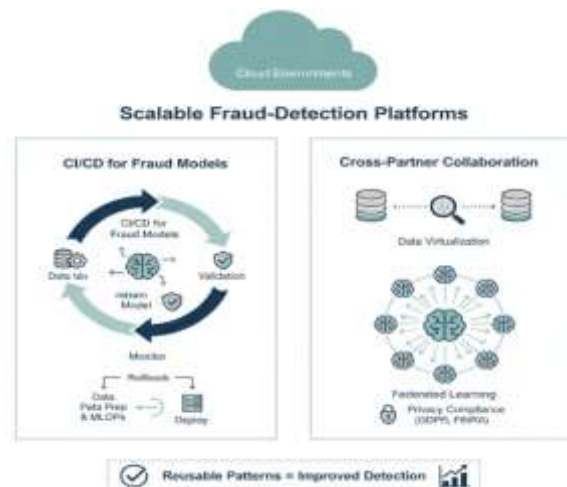


Fig 4: Scalable Privacy-Preserving Fraud Detection: A Cloud-Native Architecture Integrating CI/CD, Data Virtualization, and Federated Learning

2. Data Virtualization and Federated Learning: Model training may need to be based on data from multiple partners, affected by privacy regulations such as GDPR and FINRA publishing guidance on surveillance for broker-dealers handling customers' private data. Virtualisation



provides an attractive solution for federating access to new datasets sought for cross-partner model improvement. Additionally, federated learning captures the mosquito-and-swarm nature of fraud detection, maintaining a shared model amongst multiple partners while sharing only model weights rather than private data. This capability therefore facilitates improving model performance over shared domains while remaining compliant with privacy regulations.

7.1. Continuous Integration and Delivery for Fraud Models

Fraud is a significant issue for InsurTech companies. Criminals target the growing flow of digital data, assets, and money in the sector. Fraud-control measures must detect fraudulent activities while minimizing the disruption to genuine customers. Agentic AI can assist with fraud-detection tasks and adopt a cloud-native microservices architecture to aid responsiveness, scalability, and resilience. Pattern and model definitions help guide the practical development of adaptive fraud-detection platforms.

To reduce latency between fraud-pattern detection and response, the related models must continuously capture evolving crime patterns. Continuous integration and continuous delivery (CI/CD) pipelines automate this process for a wide range of fraud-detection models. Failures of distributed fraud-detection patterns and model cycles require rapid recovery from previous working versions. Fraud-detection models also need constant monitoring of their detecting power and appropriateness for the context to stabilize risks. Adaptation of these models must thus be automated. During these adaptation cycles, the lower risk of relapsing into previous stable states makes a major difference. This ratio helps both the estimation and presentation of relapsing risk from an adaptive model or pipeline.

7.2. Data Virtualization and Federated Learning

A fundamental precondition for reliably training any machine-learning model is access to relevant and representative input data, which may be challenging for

fraud models: congregation might not be feasible due to privacy concerns, much of the suitable information is under the exclusive ownership of individual partners, critical samples are highly imbalanced, and the need for continual adaptation to evolving behaviors makes retaining old data less relevant. Data virtualization provides a solution for accessing suitable information, either by exclusively consuming it from partners so that no physical copies are transmitted or by maintaining a collective feature repository in a read-only storage with input strata assigned to the respective partner owner.

In cases where partners collect complementary and private information, federated-learning techniques allow using these datasets for jointly training improved models while keeping the original information decentralized. For detections present in only one partner's domain or where data spanning the whole risk space is retained by a specific partner, direct collaboration is unnecessary, and any cross-domain learning must happen at a higher level of abstraction. The model must therefore be capable of externalizing the concept that it exploited while moving across domains to facilitate subsequent detection and be complemented with incoming domain-specific knowledge sources to achieve a decision-making base free from fatal knowledge lacunae.

8. Conclusions

Results support the viability of agentic AI in InsurTech fraud detection. Several approaches to creating agentic capabilities within detection-task workflows have been identified. Moreover, the underlying cloud-native architecture facilitates the deployment of scalable detection models on large datasets, enabling real-time prediction and risk scoring for transactions within fraud-prone domains. Key contributions include the following:

Improved detection performance for specific fraud types through calibration using risk scores or more abusive cohort profiles; use of lighter, less comprehensive detection models in risk-prone domains; temporal modelling for detection tasks with changing baseline behaviour;



construction of blight detection models on legislation-compliant, easily accessible data. However, these contributions are limited to the data domains available; data-access challenges continue to hamper advance across domains and jurisdictions.

Research avenues for advancing agentic AI in fraud detection are summarised as follows. First, methods should be developed to more clearly define fraud activity and typical profiles based on detection-/domain-specific criteria. Such profiles could then serve as detection-model calibration targets for specific fraud types and risk-prone domains or as guides for analysing drift within fraud-mitigation cohorts. Second, since behaviour over time is integral to risk assessment, continuous monitoring of established normal baseline behaviours, deviations from such baselines, and the emergence of new risk types are warranted. Third, capturing, analysing, and modelling novel fraud types and associated behaviours is essential to support the service, insurance, and banking sectors that together comprise the rapidly evolving FinTech landscape.

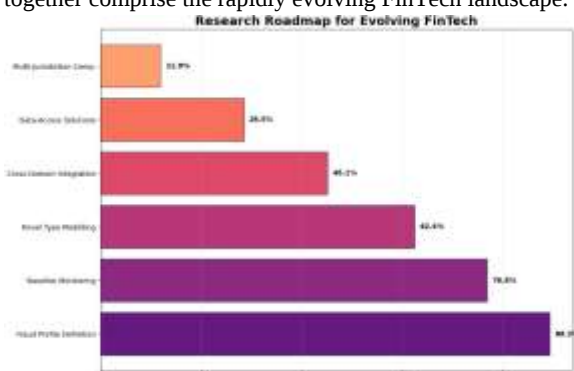


Fig 5: Research Roadmap for Evolving FinTech

8.1. Future Directions and Implications for Agentic AI in Fraud Detection

Agentic AI can contribute to the detection and mitigation of fraud in InsurTech systems by operating either passively, by delivering warning signals that can be translated into manual actions, or actively, by assuming responsibility for the execution of disruptive actions such as blacklisting accounts or restricting transactions in real time. Levels of

decision-making autonomy, as well as the corresponding delegation of responsibility and accountability, can therefore greatly vary and must be defined according to the operational context and the regulatory constraints of each specific use case. Indeed, depending on the risk associated with the action being undertaken and the potential damage due to either an incorrect decision or a data breach, some kind of human-in-the-loop mechanism may be required, complemented by auditing procedures and compliance-related controls.

To facilitate the development of practical cloud-native agentic applications for fraud detection, a set of design patterns have also been articulated. These reusable construction templates provide guidance for implementing Platform-as-a-Service environments capable of scaling the microservice architectural style beyond the isolated needs of individual applications, allowing the creation of dedicated fraud-detection platforms that fulfil the requirements of a community of partners but are not bound to any of Dyer's master Data rules. When adopted in combination with continuous integration and delivery pipelines tailored to fraud models, data virtualization strategies that address the unique data-access constraints of such use cases, and federated-learning techniques that comply with privacy regulations, they offer a complete set of recommendations for automatically detecting deviant behaviours at both the individual and collective level. Despite being limited to the domain of fraud detection, the research provides a roadmap for extending the use of Cloud-Native Agentic AI to a wider range of use cases that encompass proactive adaptive mitigation of detected threats. By establishing forms of cooperation, sharing, negotiation, and appeal among agents, not only can parallelism and thus performance be enhanced, but advanced architectures can be defined that satisfy the objectives outlined for Cloud-Native Agentic AI in a wider selection of settings. Beyond the immediate technical considerations, the evolution of fraud models towards an agentic form also raises relevant data-science and data-engineering challenges, especially when addressing the task of building real-time scoring models. Cloud-Native Agentic AI thus offers a concrete foundation for operationalizing



the concept of agentarity in AI used for real-time risk management and mitigation.

9. References

- 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.
- Farbmacher, H., Löw, L., & Spindler, M. (2022). An explainable attention network for fraud detection in claims management. *Journal of Econometrics*, 228(2), 244–258.
- Gottimukkala, V. R. R. (2020). Energy-Efficient Design Patterns for Large-Scale Banking Applications Deployed on AWS Cloud. *International Journal of Advanced Research in Computer and Communication Engineering (IJARCCE)*, DOI, 10.
- Severino, M. K., & Peng, Y. (2021). Machine learning algorithms for fraud prediction in property insurance: Empirical evidence using real-world microdata. *Machine Learning with Applications*, 5, 100074.
- Gomes, C., Jin, Z., & Yang, H. (2021). Insurance fraud detection with unsupervised deep learning. *Journal of Risk and Insurance*, 88(3), 591–624.
- Şahin, Ö. E., Ayvaz, S., & Çalınfidan, E. (2020). Comparison of machine learning models for predicting fraudulent claims in the insurance sector. *Journal of Information Technologies*, 13(4), 479–489.
- 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.
- 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.
- Aitha, A. R. (2021). Dev Ops Driven Digital Transformation: Accelerating Innovation In The Insurance Industry. *Journal of International Crisis and Risk Communication Research*.
- Zhu, X., Ao, X., Qin, Z., Chang, Y., Liu, Y., He, Q., & Li, J. (2021). Intelligent financial fraud detection practices in post-pandemic era. *The Innovation*, 2(4), 100176.
- Liu, Y., Ao, X., Qin, Z., Chi, J., Feng, J., Yang, H., & He, Q. (2021). Pick and choose: A GNN-based imbalanced learning approach for fraud detection. *Proceedings of the Web Conference 2021*, 3168–3177.
- Benchaji, I., Douzi, S., El Ouahidi, B., & Jaafari, J. (2021). Enhanced credit card fraud detection based on attention mechanism and LSTM deep model. *Journal of Big Data*, 8, 151.
- Cheng, D., Xiang, S., Shang, C., Zhang, Y., Yang, F., & Zhang, L. (2020). Spatio-temporal attention-based neural network for credit card fraud detection. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(1), 362–369.
- Inala, R. (2022). Cross-Domain MDM Integration Using AI-Driven Data Governance: A Case Study In Financial Technology Architecture. *Migration Letters*, 19(2), 280-304.
- Craja, P., Kim, A., & Lessmann, S. (2020). Deep learning for detecting financial statement fraud. *Decision Support Systems*, 139, 113421.
- Błaszczyński, J., de Almeida Filho, A. T., Matuszyk, A., Szeląg, M., & Słowiński, R. (2021). Auto loan fraud detection using dominance-based rough set approach versus machine learning methods. *Expert Systems with Applications*, 163, 113740.
- Eckert, C., & Osterrieder, K. (2020). How digitalization affects insurance companies: Overview and use cases of digital technologies. *Zeitschrift für die gesamte Versicherungswissenschaft*, 109, 333–360.
- Owens, E., Sheehan, B., Mullins, M., Cunneen, M., Ressel, J., & Castignani, G. (2022). Explainable artificial intelligence (XAI) in insurance. *Risks*, 10(12), 230.
- Butt, U. A., Mehmood, M., Shah, S. B. H., Amin, R., Shaukat, M. W., Raza, S. M., Suh, D. Y., & Piran, M.



- J. (2020). A review of machine learning algorithms for cloud computing security. *Electronics*, 9(9), 1379.
20. Inala, R. Designing Scalable Technology Architectures for Customer Data in Group Insurance and Investment Platforms.
21. Shamshirband, S., Fathi, M., Chronopoulos, A. T., Montieri, A., Palumbo, F., & Pescapè, A. (2020). Computational intelligence intrusion detection techniques in mobile cloud computing environments: Review, taxonomy, and open research issues. *Journal of Information Security and Applications*, 55, 102582.
22. Pölöskei, I. (2021). MLOps approach in the cloud-native data pipeline design. *Acta Technica Jaurinensis*, 15(1), 1–6.
23. Ruf, P., Madan, M., Reich, C., & Ould-Abdeslam, D. (2021). Demystifying MLOps and presenting a recipe for the selection of open-source tools. *Applied Sciences*, 11(19), 8861.
24. Liu, Y., Ling, Z., Huo, B., Wang, B., Liu, Y., & others. (2020). Building a platform for machine learning operations from open source frameworks. *IFAC-PapersOnLine*, 53(5), 704–709.
25. Tamburri, D. A. (2020). Sustainable MLOps—Trends and challenges. In *Proceedings of the 2020 22nd International Symposium on Symbolic and Numeric Algorithms for Scientific Computing (SYNASC 2020)* (pp. 17–23). IEEE.
26. Calvaresi, D., Dicente Cid, Y., Marinoni, M., Dragoni, A. F., Najjar, A., & Schumacher, M. (2021). Real-time multi-agent systems: Rationality, formal model, and empirical results. *Autonomous Agents and Multi-Agent Systems*, 35, 12.
27. Segireddy, A. R. (2020). Cloud Migration Strategies for High-Volume Financial Messaging Systems.
28. Gronauer, S., & Diepold, K. (2022). Multi-agent deep reinforcement learning: A survey. *Artificial Intelligence Review*, 55, 895–943.
29. Pesce, E., & Montana, G. (2020). Improving coordination in small-scale multi-agent deep reinforcement learning through memory-driven communication. *Machine Learning*, 109, 1727–1747.
30. Mao, H., Liu, W., Hao, J., Luo, J., Li, D., Zhang, Z., Wang, J., & Xiao, Z. (2020). Neighborhood cognition consistent multi-agent reinforcement learning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(5), 7219–7226.