

Research Article

Agentic AI for Next-Generation Insurance Platforms: Autonomous Decision-Making in Claims and Policy Servicing

Keerthi Amistapuram

Lead Software Developer

Received: 12/09/2025;

Revision: 23/09/2025;

Accepted: 04/10/2025;

Published: 19/10/2025

*Corresponding author: Keerthi Amistapuram

Abstract: Autonomy in AI can improve trust, scalability, efficiency, and responsiveness. This is particularly pertinent in claims processing and policy servicing, where labour needs are diminishing but demand peaks are increasing. An insurance technology platform that uses agentic AI to automate these activities would enhance the speed, quality, and efficiency of service delivery. If these autonomous systems are made trustworthy, their deployment would not only satisfy current labour shortages but also improve user sentiment in engagements that have traditionally been painful, frustrating, and expensive to manage. Acting on behalf of insurance companies, agentic AI would decide whether a claim should be settled or lead to escalation for further assessment. In the course of a policy life, agentic AI would address servicing requirements such as endorsements, renewals, and compliance checks, adjusting pricing in real time as new information becomes available. These transformations would benefit insurers whose investment in agentic AI remains aligned with the appropriate architectural paradigms—modular architecture focused on business goals, privacy-by-design data architecture, and decision-making frameworks that establish certifiability boundaries in low-risk domains such as claims and servicing. The links between claims processing, policy servicing, and architectural paradigms are further explored in the corresponding sections. Index Terms—Agentic Artificial Intelligence, Autonomy, Trustworthy AI, Insurance Technology, Claims Processing, Policy Servicing, Scalability, Efficiency, Responsiveness, Labour Shortages, Service Automation, Decision-Making Frameworks, Privacy-by-Design, Modular Architecture, Real-Time Pricing, Compliance Checks, User Experience, Certifiability Boundaries, Low-Risk Domains, Digital Transformation.

Keywords: Agentic Artificial Intelligence, Claims Processing, Policy Servicing, Trustworthy AI, Insurance Technology.

FOUNDATIONS AND CONTEXT

A comprehensive examination of principles for agentic AI in claims processing and policy servicing necessitates a grounding in agency within AI. Considerations of ethics, law, and regulation have particular salience for organizations deploying autonomous agents in domains demanding the exercise of judgment and decision-making. Moreover, the concentration of liability across all agentic components of an insurance technology ecosystem dictates careful articulation of the industry landscape and stakeholder roles. These foundations must be linked to the assumptions underlying the discussion of agentic architectural paradigms, particularly concerning data architecture and decision-making frameworks as addressed in the respective sections. Determination of the degree of agency required by the AI under consideration is fundamentally dependent on the task in question, with the requirement for transparency and explainability still applicable at all levels of independence. In the specific case of insurance applications, two broad classes of operation can be distinguished—claims processing and policy servicing—each composed of a series of decisions made at various levels of agency as defined herein. Identification of these decisions provides a basis for connecting these topics to the discussions of operational, social, regulatory, and ethical controls given later.

What constitutes agency in AI? How autonomous are agentic systems? What are the decision points that demarcate areas of agency? Answering these questions enables a better architectural design of insurance AI and provides the basis for subsequent sections on compliance, governance, trust, risk, and certification. Agency can be broadly defined as the capacity of an AI system to act independently on behalf of a user using its own internal model of the world. As a specific formulation, agency encompasses systems exhibiting agency in the sense of Scene-Understanding Vehicles (SUVs). SUVs detect and interpret the world, plan actions, execute those plans, and modify their internal states, but awareness extends only so far as to enable action. Consider, for example, standard driverless cars; their ability to change state or base actions on dynamic environment conditions provide sufficient grounds for considering them autonomous agents. That is, while they accommodate first-order action-selection processes, they lack internal models that embody complex scenes. Even if these definitions are accepted, significant ambiguity still surrounds agency. Various compounded control schemes—cascade, supervisory, and parallel—provide a hierarchy of control; when the higher layer issues commands, the subordinate layer possesses little or no influence. Different operational constraints can clearly alter the level of freedom of an agent. In general, different degrees of autonomy exist. For example, Tesla's Autopilot represents a low level of autonomy, requiring the driver to remain

A. Definitions of Agency in AI

Name: Keerthi Amistapuram

focused on driving; SAIC's driverless bus possesses an even more constrained set of allowed actions, whereas Waymo operates completely driverless in selected regions. Moreover, degrees of autonomy can change dynamically. As with any other vehicle, a Tesla can change state from fully manual to fully autonomous.



Fig. 1. Agentic AI in Insurance: Compliance, Auditability, and Explainability in Claims Processing

B. Ethical, Legal, and Regulatory Considerations

The implementation of agentic AI in claims processing and policy servicing must comply with ethical, legal, and regulatory requirements. Those obligations shape trust and governance structures, enhance customer experience, and mitigate potential risks. Establishing an actionable set of rules and controls requires information about the underlying obligation types and – to some degree – the roles fulfilled by insurers, tech providers, regulators, and customers. Most importantly, the agent must comply with the explicit and implicit rules protecting stakeholders and the broader community, as overseen by regulatory bodies. The implementation of claims processing and policy servicing agentic AI – without delegation of such compliance obligation to other insurance AI – therefore needs to comply with the following ethical, legal, and regulatory guidelines. Auditability of all decision-making processes allows compliance assessments and the development of corrective measures for detected non-compliance. Similarly, explainability of automated decision-making or policy changes is vital for user trust, either towards the systems themselves or in the responsible human agents. Such auditability and/or explainability requirements align with the ongoing debate around the accountability of automated decision-making processes in other fields.

C. Industry Landscape and Stakeholders

In the insurance context, agents are the key stakeholders: the insurers who offer the coverage and the tech providers who design, build, and operate the agentic AI. These

systems must also satisfy the requirements of regulators, such as the European Union Agency for Cybersecurity and the European Data Protection Board, who ensure that citizen rights are upheld, fraud is deterred, proper supervision is exercised, and any damage is remediated. They are also influenced by the end customers who avail of insurance products, whose experience, trust, and satisfaction will ultimately determine whether the services meet their needs. Such an approach to agentic AI will release workers from tedious operational tasks, allowing them to focus on complex, non-repetitive challenges and thus introduce a human touch to the emotionally charged claims process. It can also enhance the customer experience and build public trust in AI technologies, ultimately contributing to a more accessible insurance ecosystem and wider society. Insurers and regulators are the most visible stakeholders but not the only ones that need to be classified. The agentic AI design choices, workflows, governance structures, and operational controls also need to be examined for their implications on the business and workforce model, the customer experience, and the impact on wider society. Many of these topics recur in the relevant sections across the agentic AI service design and operational audit domains. Specific aspects are summarized and cross-referenced with the relevant decision points throughout the claims administration process.

ARCHITECTURAL PARADIGMS FOR AGENTIC INSURANCE AI

Agentic AI within next-generation insurance platforms for claims processing and policy servicing demands a finely tuned architectural approach. Core architectural assumptions include

(1) modular, task-specific AI agents perform individual functions and collaborate with one another using established orchestration patterns; (2) agentic implementations handle functions requiring agentic AI, while non-agentic functions remain supported by conventional AI; (3) modelling of customer-facing interactions places priority on establishing customer trust through preparation of a true-to-life persona backed by policy-relevant controls and disclosures; and (4) data privacy and protection is a priority and embodied using privacy-by-design principles. The use of agentic AI will especially benefit functions that lend themselves to agentic capabilities for autonomous decision-making, such as autonomous claims processing and autonomous policy servicing. When considering claims processing and policy servicing from a principled privacy and protection standpoint, the fundamental requirements can be distilled into three considerations: these tasks require substantial interactions with sensitive data; the AI-based decision-making must be explainable and accountable; and the solutions must protect customer data while satisfying regulatory rules. Supporting functions do not introduce sensitive operational data at risk or generate agentic AI-directed activities, and the core requirements do not require full privacy controls. These observations shape decision-making around the use of agentic AI within autonomous claims processing and autonomous policy servicing.

A. Modular AI Agents and Orchestration

Agent-based systems are inherently modular, as they often perform largely autonomous roles and may come from different sources or even subsist in different technical environments (e.g. coding languages, software frameworks or cloud versus on-premise solutions). It is vital to ensure that the language and protocols used for agent communication and cooperation allow for seamless interaction among agents that belong to different modules. Such interaction is possible through the use of standard, widely known and pervasive (e.g. IP) languages and protocols, such as HTTP and JSON. A key feature of modular architectures is that failure is isolated to the agent that initiated a fault, thereby making fault management easier (e.g. by establishing an internal watchdog that monitors the execution of individual agents on a service-by-service basis). The advantages of agentic designs are realised when different agent types with different functionalities and specialisations are employed, both for redundancy and capability enhancement. A particular

advantage of deploying such agentic systems is that agents can readily be produced for almost any task, or combination of tasks, and therefore can be readily procured from the many AI development companies now offering such capabilities. This is especially important for tasks that may be once-off, rare, or for which internal development capability does not exist. Incorporating modularity, means that task-specific AI agents may be procured from vendors who provide their services on a pay-per-use basis. These stickiness-reducing and cost-reduced lemniscate and yin-yang relationships contrast with traditional models that usually flow in a linear manner, but still require that services of certain key providers be continued even when they don't have the best offer. The modularity of agentic systems can also facilitate recovery planning by allowing easy identification of parties that are critical for particular services. Modular agentic systems tend to be networks with a polycentric governance structure in which risk is shared by all parties participating in the network.

Derivations — Symbols & Notation

Equation 1 — Agentic Intelligence Core Function (decision rule)

Goal (per paper): An agent chooses an action on behalf of the insurer, balancing loss/costs and constraints (privacy, compliance, certifiability).

Bayes risk with constraints → optimal policy

We minimize expected loss under constraints:

$$\pi(a | x) = \arg \min_a \pi \in \Delta(A) \sum \pi(a | x) E \quad (1)$$

$$E_{y \sim p(y|x)}[L(a, y)] \text{ s.t. } E[g_k(a, x)] \leq 0 \quad \forall k. \quad (2) \text{ Using Lagrange multipliers } \lambda_k \geq 0, \text{ the Lagrangian is:}$$

$$L(\pi, \lambda) = \sum_a \pi(a | x) E[L(a, y)] \quad (3)$$

$$+ \sum_k \lambda_k \sum_a \pi(a | x) E[g_k(a, x)]. \quad (4)$$

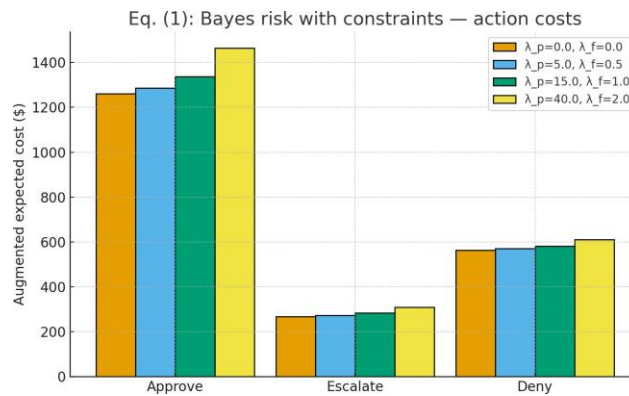


Fig. 2. Bayes risk with constraints action costs

Because the objective is linear in $\pi(a | x)$ over the probability simplex, the optimum places mass on actions minimizing the augmented cost:

$$a(x) = \arg a \in A \min E[L(a, y)] + \sum_k \lambda_k E[g_k(a, x)]. \quad (5)$$

B. Data Architecture and Privacy-by-Design

Data lineage from data generation to consumption must be clearly traceable, and any personal or sensitive information must be

clearly labeled. Failure to do so can inhibit reuse, increase the risk of inadvertent disclosure, and cause e- discovery and privacy-compliance difficulties. Sufficient information must be retained to enable customers and regulators to assess and understand any decision made on their behalf. Hence, it must be easy to apply privacy-by-design principles, especially data minimization and privacy preservation, without major engineering effort. This includes implementing privacy controls and mechanisms before the point of execution. More broadly, data access must be governed through policies and mechanisms that ensure compliance with legal and regulatory requirements, the data subject’s instructions or consent, and any relevant internal standards. Interoperability requirements stemming from data security, safety, risk governance, privacy, and other considerations must be articulated together with appropriate data standards. The service that processes the data must have an explicit and sufficiently restricted view of the data, sharing with other services information that is relevant to those services while safeguarding security, privacy, and regulatory requirements. Satisfaction of such considerations contributes directly to the robustness of agentic AI systems and establishes the foundation for higher quality answers. Agentic AI service quality critically depends on realism and reliability. As with all AI models, continuous monitoring of performance, supported by an appropriate accreditation model, is necessary. This includes the generation of tests and controls robust enough to catch even adversarial inputs. When key components perform unsatisfactorily, a model-failure mechanism

λ priv	λ fair	Cost(Approve)	Cost(Escalate)	Cost(Deny)	Chosen a^*
0.0	0.0	1260.0	267.0	564.0	Escalate
5.0	0.5	1286.0	272.5	570.5	Escalate
15.0	1.0	1337.0	283.0	582.0	Escalate
40.0	2.0	1464.0	309.0	610.0	Escalate

TABLE I AUGMENTED COSTS POLICY

Takes over. In highly sensitive areas, such as fraud detection, it is advisable to maintain multiple alternatives within the overall framework. The systems in place must make the effort to genuinely reduce false positives and negatives without compromising business requirements. Quality, redundancies, resilience, safety margins, and the unavoidable need for human operating staff must all be monitored.

C. Decision-Making Frameworks and Certifiability

Risk boundaries are defined for automated AI decision-making processes. When decisions exceed these boundaries, it is expected that human intervention will be required at the next applicable check-point. Such capacities can be described according to levels of agency: organisational “superintelligence” lies at the extreme end of the spectrum, exercising ultimate control over data selection and machine learning for the entire insurance ecosystem. The lowest level of agency, defined as organisational non-superintelligence, fails to maintain controllable data selection and unregulated machine learning. A distinction is made in terms of decision-certifiability: decisions that are certified by appropriate authorities or an equivalent algorithm can be reasonably executed autonomously, while those that are non-certified need human intervention in order to comply with applicable regulations and organisational governance and modelling. As such, decisions associated with insufficient organisational-user-quality labelling, risk scores, data-privacy ratings or analogous quality point systems should not be implemented autonomously until appropriate resale-to-riskwatch mechanisms are in place, nor should decisions require review of a human user on the solid reputation of the commendable insurance products or services available from the organisation requesting the user-layer evaluation.

AUTONOMOUS CLAIMS PROCESSING

An arrangement of agentic processing tailored to the claims lifecycle, from intelligent fraud detection to settlement execution, is presented. The approach encompasses the detection of events indicative of fraudulent behavior, the triaging of new claims, and customer-interaction protocols around claims processing. Considerations such as fraud-related models, explanation and auditability requirements, customer interaction policies, disclosure rules, and settlement pathways are highlighted. Various signals can indicate increased likelihood of fraudulent claim submission. Anomalies in claim

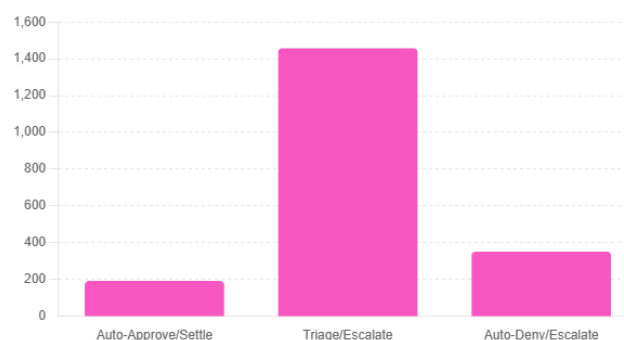


Fig. 3. Triage Outcomes by Thresholds

Major events or abnormal weather patterns) also offer valuable clues. Signals may be useful individually or in combination. When captured, they can operate as a detection layer around incoming claims and can enable focused investigation along a subset of incoming claims (and possibly claimants) before subsequent approval. Understood in a causative way, these signals possess clear interpretable axes for when flagged by the model, making future explanation to both regulators and consumers straightforward. Audit trails around model predictions provide a clear record of evidence leading to escalation of claims into fraud investigations and decisions made there, facilitating internal governance procedures. Insurance operations typically follow an operational triage process around new claim notifications or requests. Similar modeling and escalation patterns could easily be developed for validating claim repairs. An arrangement enabling the automatic settlement of claims against policies with simple cash payouts is presented. When total claim amounts across policies are themselves cash payouts, settlement may be truly automatic, freeing consumers from even entering banks to start the process.

Equation 2 — Claims Automation Model (fraud score + triage)

Goal (per paper): Detect fraud signals, route to triage, and auto-settle/deny based on calibrated confidence/coverage.
Step 1: Calibrated fraud probability

Let a linear predictor $z = wTx + b$. The standard calibrated model:

Amounts and claim patterns often provide useful indication, as do inconsistent signals in their accounts across claims and on social media networks. Additional signals pointing at fraud

$$pfraud(x) = \sigma(z) = \frac{1}{1 + e^{-z}}$$

Step 2: Threshold triage with coverage logic detection in the broader operating environment (e.g., global Pick two thresholds $0 < \tau_{low} < \tau_{high} < 1$ and a coverage

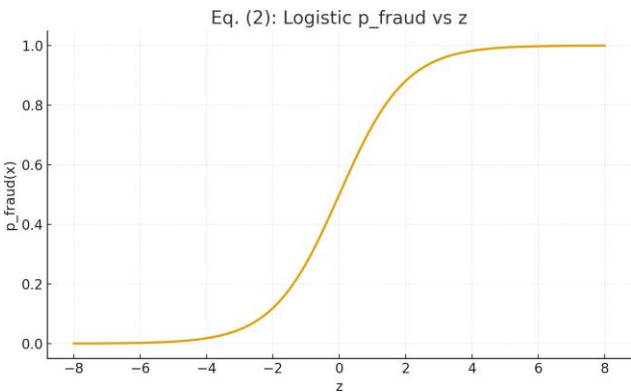


Fig. 4. Logistic p fraud vs z

Outcome	Count
Auto-approve	529
Auto-deny/Escalate	860
Triage/Escalate	1611

TABLE II Triage counts indicator $C(x) \in \{0, 1\}$. Decision rule:

Tioning or data retrieval if required. Completion of the triage’s validation step without explicit falsification of the fraud flag is interpreted as providing some support of the fraud signal on the overall claim. The explicit results of the fraud-detection task and their explanations are nevertheless maintained in an easily accessible deck for auditing and data-usage purposes. The complete deck, combined with guidance and supervision from the model-governance provision (potentially through an external, duly qualified body), allows the insurer to remain compliant with explainable-output regulations at least at that level while gradually working toward minimizing the signals from the fraud-detection toolset.

B. Automated Triage, Validation, and Settlement

After initial fraud detection and risk assessment, triage is automated based on detection confidence and policy coverage and conditions. Claims that are assessed to be low risk and fall within defined policy terms proceed to validation checks, such as automated verification of evidence requirements. Highly confident detections lead directly to settlement along preap- proved pathways, with information automatically shared with the claimant. In most claims environments, three factors drive the need for triaging claims: variability in fraudocclusion capabilities, controlled remediation costs, and the desire to remediate only the

most egregious cases. The vendor-supplied checklists of requirements in the cloud or premises-based solution handle this aspect efficiently. Detected noncompliances or audits at lower confidence levels then trigger normal-scale processing of the cases involved. This provides for intelligent fraud detection at scale, with information from the triaging

Step 3: Expected risk control

Choose τ 's to cap expected fraud loss:

- $E[L(\delta(x), y)] \leq B \Rightarrow \text{calibrate } (\tau_{\text{low}}, \tau_{\text{high}}) \text{ on validation data s.t. the bound holds. Visuals produced:}$
- Logistic curve $p_{\text{fraud}}(x)$ vs. z (lineplot).
- Bar chart of simulated triage counts under
- $(\tau_{\text{low}}, \tau_{\text{high}}) = (0.2, 0.7)$ $(\tau_{\text{low}}, \tau_{\text{high}}) = (0.2, 0.7)$.

A. Intelligent Fraud Detection and Risk Assessment

Fraud detection operates in the early stages of the claims process, analysing incoming claims and, where applicable, flagging indications of fraudulent activity for further examination. In most incidents, fraud detection is not a black-and-white matter, and its objective becomes a risk score that indicates how likely a claim may be fraudulent. This stage is idempotent—multiple solutions flagging identical claims for further investigation do not change this status. The main output is thus merely an explanation of why the flag appeared for this unique claim. Such transparency reduces the risk of these flags being ignored or dismissed as inconsequential. The signals generated by the fraud-detection process must be legible to a human viewer, allowing for a fast, informed, and reasonable decision as to the fraud claim's validity. The flagged claims are routed, along with the signals, into the triage process for more thorough examination, which includes further manual question and validation process feeding back into updates of the fraud detection capability and policy terms structuring. The portal-style completion of settlements for low-risk claims provides key claim-tuning satisfaction factors for customers.

C. Customer Interaction and Transparency

For all questions and decisions not explicitly captured by existing models and rules of the insurance organization, interactions with customers and/or other external stakeholders remain vitally important. Hence, customer interaction and transparency of internal decisions must be clearly defined. To enhance customer trust and confidence, an appropriate communication style should be chosen based on the customer's preferences, either personal or impersonal. The technology should be able to disclose aspects of its operations and decisions in a way that is meaningful to customers, explaining and justifying decisions or providing information in easily understandable language when required. The appropriate level of explainability for a given decision is linked to the customer's familiarity with the technology, with explanations becoming increasingly comprehensible as their experience deepens. Appropriate controls for customer interaction need to be defined. Primarily, a preferred contact channel should be specified. Additionally, customers may want to restrict the type of information shared with the technology: many would desire the ability to notify the technology of information and events not organically detectable (for example, purchases being made in other geographical areas), while some may wish to inhibit communication of certain actions being undertaken (for example, a person with psychiatric illnesses may wish not to reveal to the insurer that a therapist or medical professional has been consulted again).

AUTONOMOUS POLICY SERVICING

Supporting activities catering to all phases of the policy lifecycle further enhance efficiency, customer experience, and compliance. Operations such as policy issuance, endorsements, renewals, and servicing which track adherence to the policy terms are essential. Policies can be automatically issued or renewed based on intelligent predictive models—unless triggered otherwise by an external party. Incoming endorsements are processed without delay if possible. Monitoring systems reinforce compliance with the country of domicile and other exclusions. In addition to issuing and renewing insurance covers, pricing adaptations serving customer objectives and real-time underwriting during the contract period complete the servicing space. Lifecycle activities can assign customer sign-offs and approvals, thereby improving accessibility and increasing responsiveness. These can be automated in less-exposed lines such as motor insurance, or processed with predictive decision-support systems to ensure proper appropriateness without bias.

A. Policy Issuance, Endorsements, and Renewals

Agents can automatically issue new policies when conditions are satisfied—e.g., the quoted risk premium is within a specified range, and the underwriter has given prior consent to similar risks—or issue endorsements to existing policies, generally without further validation. Auto-renewal of expiring policies takes place unless flagged by underwriters or other agents or rejected by customers. Information from past interactions helps retain customers for policy renewals and premiums in appetite for various underwriting authorities; both usually involve minimal customer engagement. Such decisions must be certifiably safe within established risk boundaries. Clear triggers and closing conditions for issuing endorsements help streamline data preparation for risk-logic engines. Risk parameters and other factors should furthermore be monitored as part of good practice lifecycle risk controls; red flags warrant additional scrutiny before adjustment. Where combining and unifying certain types of data artefacts into the final decision is critical, use-cases that lend themselves to formal test-data generation based on equilibrium and other state-based models are particularly relevant. The safety of sentient AI functioning across the full licensing spectrum demands more than just common-sense safety certifications of individual actions; almost any feature can have catastrophic real-world consequences on accident-prone areas or high-frequency negatives across many otherwise fine decisions. Safety checks,

monitoring, and backup procedures for pre- dictive maintenance and other purpose-specific capabilities therefore require particular care to ensure back-fill options

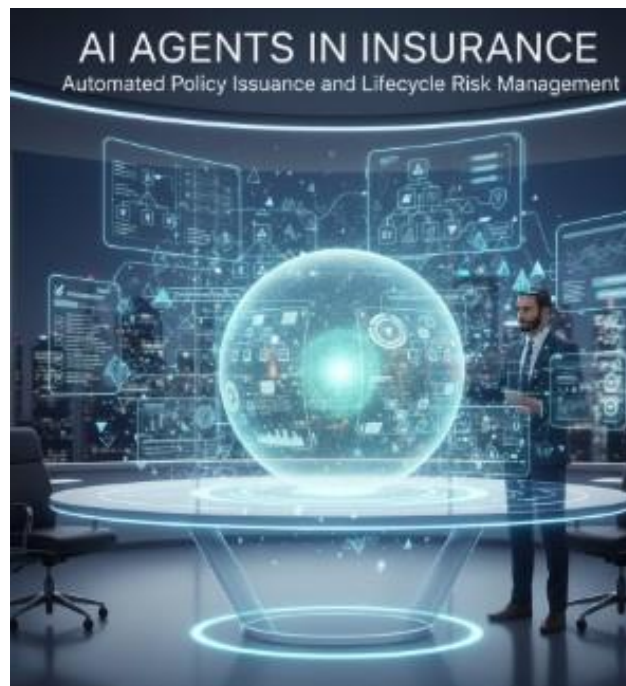


Fig. 5. AI Agents in Insurance: Automated Policy Issuance and Lifecycle Risk Management

Exist and are obviously flagged for scrutiny as AI decision- making fatigue sets in far beyond human thresholds.

B. Pricing Adaptation and Real-Time Underwriting

Dynamic pricing and real-time underwriting in insurance require regular re-evaluation of pricing models and exposure risk assessments. For successful integration, insurers must build advanced granular analytical capabilities, create internal systems that act as new-age data platforms, develop new and better algorithms, and ensure real-time data flows to pricing systems. Underlying these capabilities, the data feeding the pricing models must be of sufficient quality—accurate, complete, timely, relevant, and consistent. Moreover, underlying data sources must support integration of heterogeneous data from different environments. Data partners must be defined based on credit attributes. In regions where a large portion of the population is not banked, insurers can either partner or rely on Information-as-a-Service (IaaS) players to provide data to identify risk exposures. Pricing inputs must take into account risk scoring across product lines—where the score can indicate a classification or indication of bad underwriting quality—real-time capability to assess underwriting exposure, and regulatory constraints on price adjustments. Moreover, the requirement must be implemented in a way that allows rules to be defined without software development but using a rules engine.

C. Lifecycle Servicing and Proactive Compliance

Lifecycle activities, including monitoring for intermediate events, managing endorsements, handling renewals, and undertaking periodic due diligence for compliance, can be performed autonomously. If requested, agents should ensure necessary information is requested, that policies remain compliant, and that customers remain informed. Triage procedures can manage the flow of information for more complex areas of customer interaction, such as claims submission. All customer-facing communication should be attentive to people's level of perception capabilities—whether of select group such as seniors, or the larger society. In respect of compliance, KYC procedures can be iteratively enforced as new data become available; noncompliance can trigger continuous coverage warnings or even cancellation. Business endorsements can also be added on a need basis. Further auto-triggered communications may include advice on identified occupation changes, and request for updated information according to changes identified in active monitoring. An open-ended reserve for ad-hoc requests may also be included. Proactive service ensures that users feel the company is taking interest in them from time to time and that the customer journey is truly self-servicing. Thereby, the overall customer experience becomes a lot friendlier and user confidence in the product and service likely increases, notwithstanding the fears of machine-led decisions replacing human contact.

GOVERNANCE, TRUST, AND RISK MANAGEMENT

Governance structures, accountability models, and risk controls are presented. The discussion weaves cross-references to sections on explainability, certifiability, and operational risk. Society has entrusted the insurance sector with providing a safety net against financial loss arising from various risks, and individuals expect compensation without delay when these events occur. To manage this risk transfer creation, insurers have set up massive teams, invested heavily in sophisticated crime detection methods, and walked with a large stick ever since Regulation 1 was born. Therefore, it is crucial for any agentic AI

system being established for next-generation insurance platforms to create even greater trust and transparency among the users of the systems to ensure smooth transitioning. Any insurance organization that deploys agentic AI must be Governance of such advanced systems would require external bodies to audit these systems periodically, and models would need to be versioned much like financial systems. Audit cycles would need to be defined along with external reviews by trusted AI organizations. Audit logs, ensure impact assessments during the design phase, and justifications for why the creation of such models is needed must be maintained in future releases. Also, it is important that for any decisions made where the primary stakeholders are the customers, and such decisions are made explainable as well so proper refunds issued by the system or changes made can be easily understood by the customer.

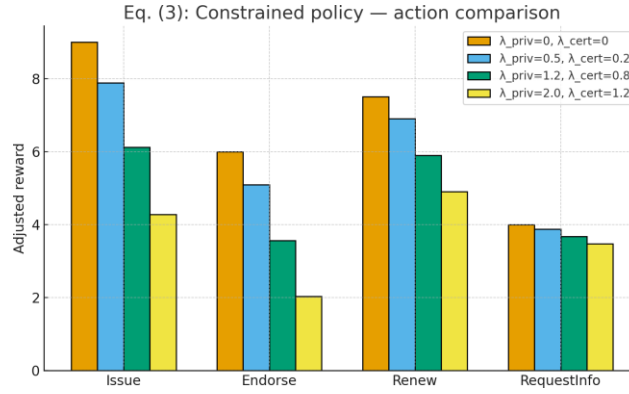


Fig. 6. Constrained policy action comparison

Equation 3 — Policy Servicing Optimization (constrained MDP)

Goal (per paper): Optimize issuance/endorsements/renewals with certifiable safety and cost/risk tradeoffs.

Step 1: MDP formulation

State s (policy + customer context), action a (issue, endorse, renew, request info), transition $P(s'|s, a)$, reward $R(s, a)$ (e.g., margin – cost), discount $0 < \gamma < 1$.

Unconstrained Bellman optimality:

- $V(s) = \text{amax}[R(s, a) + \gamma E s'[V(s')]]$. (8)
- Step 2: Add compliance/risk constraints
- Let $c_j(s, a)$ be risk/compliance costs (e.g., fairness, privacy, certifiability). Constrained objective:
- $\pi_{\text{max}} E \pi[t = 0 \sum \omega_j R(st, at)] s.t. E \pi[t = 0 \sum \omega_j c_j(st, at)] \leq d_j$.
- (9)
- Step 3: Lagrangian relaxation $\rightarrow$ solvable Bellman
- $R(s, a) = R(s, a) - \sum_j \lambda_j c_j(s, a), \lambda_j \geq 0$, (10)
- $V^*(s) = \text{argmax}_a [R(s, a) + \gamma E[V^*(s')]]$. (11)

Dual ascent on λ enforces constraints; resulting policy is certifiably safe within set budgets d_j .

Visual produced:

Efficiency vs automation (links Equation 3's value improvements with Equation 6's efficiency metric).

A. Model Governance and Auditability

To attend to ethical, legal, and regulatory compliance, agentic AI in claims and policy servicing requires appropriate model governance. Audit cycles should be defined, versioning processes formalized, and dedicated external reviews conducted at appropriate intervals. Evidence of such governance is fundamental for an insurance AI to gain user trust. Periodic review should cover important model attributes such as sufficient trustworthiness, risk certifiability (per Section 2.3), robust defenses against adversarial inputs (per Section 6.2), explainability (per Section 5.2), fairness (per Section 7.3), and security (per Section 6.3). Audit cycles should thus ideally align with these other assessments, enabling interconnected review via a read-focused approach. Wider industry adoption would benefit from structured certification schemes that respond to regulatory requirements for model auditability and approval. Development teams would therefore likely undertake audits in anticipation of external validation by assurance providers and regulators. Such auditing could also support parallel pathways for informal assurance, with external reviews conducted at lower frequency on the model-as-a-service path and for less trusted customers. Independent validation by other user organizations would be appropriate if internal tools were to be made available and if validation demanded by any user periodically warranted involvement of external reviewers.

B. Explainability, User Trust, and Accountability

Every decision and disclosure that affects customers or is likely to influence their trust must be explainable in a manner that is appropriate for the intended audience and the context of the decision. For internal decisions, customers should be provided with explanations only when the decision is non-trivial and when the absence of an explanation would hinder proper validation of the decision (e.g., for very high-risk decisions). When requests for explanations are made by customers, clear and understandable explanations must be supplied promptly and free of charge. Many insurance transactions are inherently complex, making it impossible to provide customers with satisfactory explanations for all of the factors that went into the decisions that affect them. Therefore, any lack of such explanations must be compensated by other trust-building measures. Examples of such measures include the timely completion of the customer journey and swift and smooth communications with customers. Communication with customers must also follow the principles of "clear and simple" communications, especially for disclosures related to customers's rights and obligations and for legal documents. Where there are doubts about the sufficiency of trust, conditions must be applied to promote customers's acceptance of the insurance offer and to enable admission to insurance on terms that do not expose insurers or other policyholders to excessive risk.

C. Operational Risk, Safety Margins, and Failover

To ensure operational continuity, services must be monitored, redundant pathways activated when necessary, and management alerted for manual intervention. These safety margins enable graceful degradation in response to predictable failures, while auditing and testing activities help drive down the incidence of unexpected faults. The transparent flow of funds and data across insulated trust boundaries informs the level of scrutiny needed for model decisions. Sufficient safety margins, however, also permit some testing of interfaces with the broader environment; external conditions can fail without causing model failure, allowing observations that support

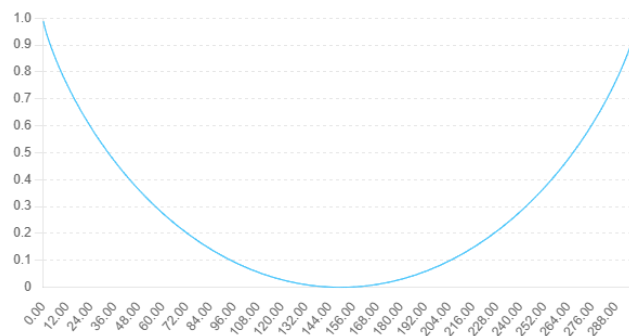


Fig. 7. Confidence from Predictive Entropy

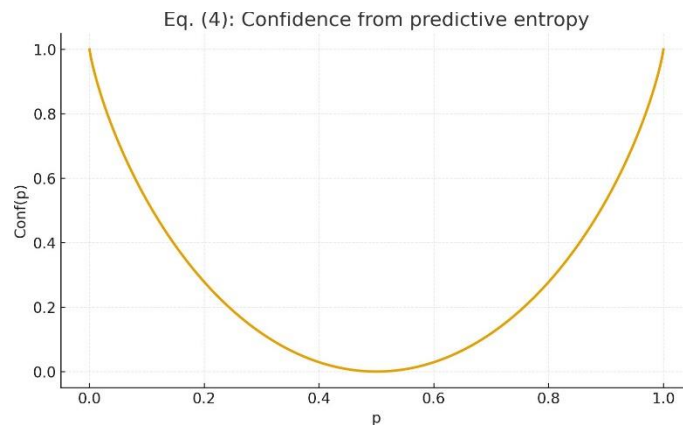


Fig. 8. Confidence from predictive entropy

Environment-hardening efforts.

- Equation 4 — Autonomous Decision Confidence (entropy-based)
- Goal (per paper): Use explainable confidence for escalate/approve decisions.
- Step 1: Predictive entropy for binary decisions
- $H(p) = -(p \ln p + (1 - p) \ln(1 - p))$. (12)
- Step 2: Normalized confidence score
- Max binary entropy is $\ln 2$ at $p = 0.5$. Define
- $\text{Conf}(p) = 1 - \ln 2 H(p) \in [0, 1]$. (13)
- Low confidence near 0.5 triggers escalation; high confidence near 0 or 1 enables automation.
- Visual produced:
- Line plot of $\text{Conf}(p)$ vs p , showing the "U" shape (lowest at 0.5).

V. TECHNICAL CHALLENGES AND MITIGATION STRATEGIES

Agentic agents remain at the forefront of advancing claims processing and automated policy servicing. Each development milestone poses specific technical challenges that warrant

p	Entropy H(p)	Confidence Conf(p)
1e-06	1.4815510057992861e-05	0.999978625737111
0.002005004008016032	0.014458296525524869	0.9791410873029232
0.004009008016032064	0.02612752411137135	0.9623059505338177
0.006013012024048096	0.03674441571896208	0.9469890136618909
0.008017016032064127	0.046676428058074704	0.9326601487142051
0.01002102004008016	0.05609810165847494	0.9190675469340333
0.01202502404809619	0.06511224289117323	0.9060628900797467
0.01402902805611223	0.07378678296817587	0.8935481741286624
0.016033032064128257	0.08217000239857866	0.8814537450297364
0.01803703607214429	0.09029799914544004	0.8697275244306778

TABLE III ENTROPY CONFIDENCE SAMPLE

Detailed consideration. Addressing these technological hurdles requires a multi-faceted approach centered on data quality, robustness to adversarial inputs, and security of smart contracts. By implementing a thorough data-control framework, introducing data-oriented testing for adversarial resistance, and establishing rigorous security assessments, it becomes possible to harmonize technical feasibility with agentic use for claims and policy servicing. Achieving agentic AI in insurance requires high-quality data to feed comprehensive, integrated decision-making models. Interoperability across domains facilitates data reuse and avoids data silos within the agents. Accordingly, the requirements for training data, decision inputs, and operational monitoring must be expressed as a coherent, multi-faceted data-control framework. Data standards ensure quality for external-sourcing decisions and prevent copy-paste transfers of harmful models. A blend of cleansing procedures leverages human expertise effectively while minimizing costs; and unified data pipelines provide a stable foundation for model-training needs.

A. Data Quality and Interoperability

Achieving high-quality data is critical for the probabilistic inferences required in agency and autonomous decision-making. Data quality encompasses accuracy, consistency, completeness, and timeliness; indicators of data fitness for use in specific applications must be defined and evaluated at key stages. To support effective risk scoring and real-time correlation of leading fraud indicators, data used for risk assessment and intelligent fraud detection must be both high quality and comprehensive, potentially leveraging alternative datasets from third parties. The quality and completeness of data classified as a common source of risk should be actively maintained. Data acquisition and cleansing strategies will be determined as part of the development and testing cycle. The ability to collect, integrate, and correlate high-quality data at the speed required for real-time pricing adaptation is also critical. Key pricing inputs such as speed, geopolitical risk, and locality sentiment are expected to be highly volatile. Risk-based pricing must therefore support rapid response times and high-quality scoping. To ensure the required data quality and availability, pricing adaptation must be integrated with other lifecycle activities and support formal certification prior to commencement. Ingesting real-time sensor data from Internet of Things devices is a promising avenue for proximal



Fig. 9. High-Quality Data in AI: Ensuring Accuracy and Real-Time Integration for Risk and Pricing

Risk-scoring evaluation. Data competent in both quality and completeness will be monitored, and active remediation steps will be taken where feasible. Supporting regulatory authorities will be consulted so that their approval is secured prior to the required real-time decision-maker readiness.

B. Robustness to Adversarial Inputs

The threat landscape for agentic AI systems encompasses various dangers, out of which malicious actions hold immense destructive potential, as seen in recent AI developments. Cybercriminals have seized upon vulnerabilities and are deploying state-of-the-art generative AI systems to launch sophisticated attacks that are easily customized. As AI technology continues to evolve, these threats will become even more intelligent, automated, and pervasive. Malicious actors will find more effective means to automate processes that were previously beyond their capabilities. According to predictions, attacks will shift from disruptive and destructive methods to more stealthy and under-the-radar channels by manipulating AI systems into helping the attacker accomplish their goals.

Unfortunately, the insurance industry is ill-equipped to combat these risks. The success of agentic AI systems, which rely on intelligent agents constantly calling one another through cloud service APIs, hinges on the inherent trust denied to humans in the industry. When commencing an interaction with another party, either human or machine, an individual must assess their level of trust and treat them with the appropriate level of caution. AI-powered systems and smart contracts can facilitate this process, allowing various players to verify that a service being rendered has completed the corresponding verification steps.

C. Security, Privacy, and Regulatory Compliance

Comprehensive threat modeling guides defence strategies against data breaches, service misuse, and exploitation by criminals—especially in rules-heavy sectors such as insurance, finance, and healthcare. Consolidating insights into a regulatory framework streamlines compliance with various local and cross-border regimes (e.g. GDPR, HIPAA, PCI-DSS, PDPO) governing data protection, consumer safety, advertising and marketing, anti-spam, anti-money laundering, anti-terrorism financing, and cryptocurrency use. The absence of a standardized compliance assessment for pre-query machine learning services (e.g. OpenAI’s ChatGPT, Google’s Gemini) allows failures to build rare elements of our digital world. Insurers should adopt systems-cross, compartmentalized defenses that offer regulators clear audit options while underpinning actual services—surveillance sensors should remain separable from money transfer capabilities. In summary, agentic AI agents for claims and policy servicing are highly sensitive systems that must operate under stringent controls, especially for development and training activities. Data privacy considerations demand tracing of sensitive data from ingestion through to query-setting results.

Equation 5 — Dynamic Risk Assessment Function (streaming/real-time)

- Goal (per paper): Risk scoring that updates with new signals (IoT, weather, social, macro).
- Step 1: Exponentially weighted update Let r_t be the instantaneous risk signal at time t . With forgetting factor $\beta \in (0, 1)$:

$$R_t = (1 - \beta)r_t + \beta R_{t-1} \quad (14)$$
- Step 2: Composition across heterogeneous sources
- With features x_t and model $p(y | x_t)$, a calibrated fraud/claim risk:

$$Risk = \alpha R_t + (1 - \alpha)p(y = 1 | x_t) \quad (15)$$
- $\alpha \in [0, 1]$ balances long-memory environment risk and current model signal.
- Step 3: Thresholds with governance buffers Operationalize with guard-bands ϵ from Equation 4’s confidence:

- auto if $\text{Conf}(\text{Risk}) \geq 1 - \varepsilon$; escalate otherwise. (16)

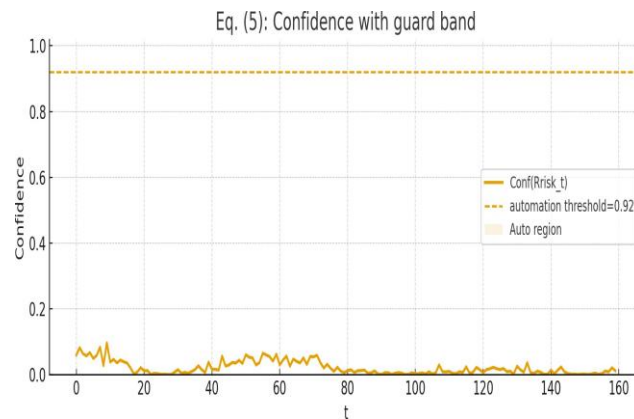


Fig. 10. Confidence with guard band

IMPACT ON LABOR, CUSTOMER EXPERIENCE, AND SOCIETY

Next-generation agentic AI directly supports workforce transformation, with semantic and other automation augmenting fundamental insurance roles: underwriting, broking, and claims adjusting. Altered tasks, however, will demand new competencies, with AI use requiring training in systems, risk management, and framework-specific guidelines. Reskilling, upskilling, or redeployment can contain overall workforce reductions, but these remain probable, unless demand growth offsets efficiency gains. Supporting change through consultation and transparent allocation of remaining work can further sustain trust, morale, and productivity. Customers can also benefit, with seamless, scalable interactions made possible by AI avoiding pernicious laziness. Ongoing, proactive compliance enables artificial non-negligence during service delivery, promoting satisfactory outcomes, desire for continued engagement, accessibility, richly explainable decisions, and ultimately, agentic AI adoption. Moreover, service or support-enhancing services can expand coverage without compromising cost or value. Nevertheless, risk-adjusted prices must remain affordable for attractive insurance products, and modelling thresholds appropriately, including for demographic bias, remains essential if unfair exclusion is to be avoided. Close-fitting wordings with personalisation, automated endorsement, and continuous fulfilment further contribute to accessible new solutions.

A. Workforce Transformation and Skill Requirements

Transformative AI technologies raise concerns about job displacement, especially for tasks susceptible to automation. Although many insurance jobs involve context-specific decision-making or interpersonal skills that are difficult to automate, the nature of human involvement may change radically. Insurance staff in claims and policy servicing functions may need to develop new skills to complement agentic AI rather than simply superseding it. Therefore, there will be demand for workforce transformation rather than wholesale job losses. Upskilling efforts should aim to familiarize staff with the augmented capabilities of AI-enabled tools and build a synergistic relationship with the technology, rather than drive a wedge between staff and systems. The need for close engagement with channel partners, such as banks, to conform to joint operating standards also necessitates a human touch that cannot be fully taken over by AI. The continued importance of qualified professionals in performing the latter stages of the customer-interaction process is also notable. In certain operational contexts, such as fraud detection, a simulation approach can be taken to further expand the bargaining power of the human-in-the-loop. Technology vendors supplying source-technology for autonomous deployments will also need to invest in new specialised skills that go beyond technical implementation. The idea is to place security, privacy, and regulatory compliance at the core of the product development life cycle right from the architecture stage, rather than as a legacy consideration that can be attended to later during the operational phase of the joint product. Close collaboration with regulatory authorities can help ensure that risk provides a safety margin in product marketing, scaling, and delivery.

B. Customer Trust, Satisfaction, and Accessibility

Considerations for trust growth encompass various aspects. Service quality is vital, particularly in sensitive matters such as insurance claim or inquiry responses. Transparency regarding AI usage heightens scrutiny. Meeting the aforementioned requirements for customer interactions supports this. Rights enforcements, such as claim denials, should be clearly communicated to foster user confidence. Adopting a responsible data utilization policy that upholds privacy and minimizes leak risks also helps. Furthermore, following the specified guidelines for explainability solidifies accountability, bolstering trust. Accessible services enable individuals with cognitive, hearing, sight, or speech limitations to communicate and execute transactions seamlessly. Adherence to WCAG 2.0 guidelines guarantees accessibility of web services, while compliance with relevant government standards ensures widespread document accessibility. Assisting in-person interactions aids customers requiring support. Making services available in multiple national and widespread languages enhances the overall experience. Such measures contribute directly to customer satisfaction and trust enhancement, ultimately establishing a more resilient and impactful overall system.

Equation 6 — Agentic Performance Efficiency (ops/econ) Goal (per paper): Tie throughput, quality, cost, and risk into

a single operational KPI for governance.

- Step 1: Define measurable components
- Throughput $T(\alpha)$: claims/policies per day as automation level
- α rises.
- Quality $X(X)Q(\alpha)$: accuracy/overturn-rate complement.
- Cost $C(\alpha)$: inclusive of compute, staff, audit. Risk multiplier $\mu(\alpha) \geq 1$; $M(\alpha) \geq 1$: inflates denominator when risk rises (from fairness, security, regulatory exposure).

Step 2: Efficiency function

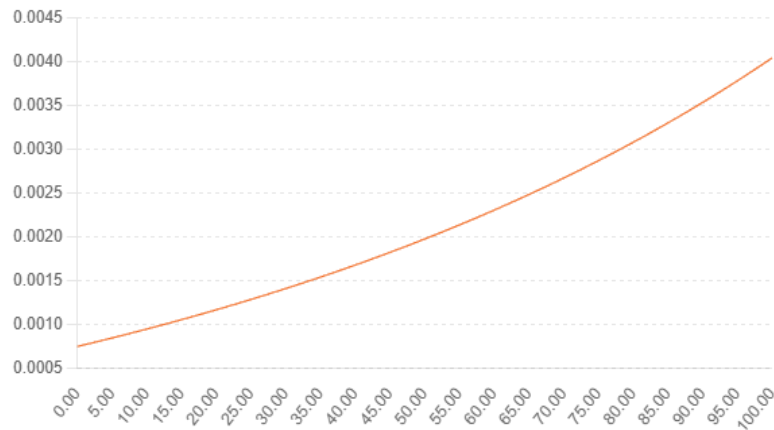


Fig. 11. Efficiency vs Automation

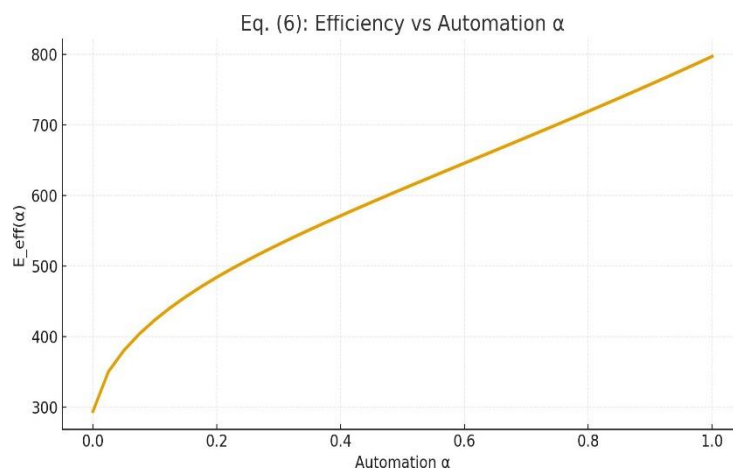


Fig. 12. Efficiency vs Automation

This scalar KPI increases with more processing and accuracy, but penalizes cost and risk.

Visual produced:

Line plot: $E_{eff}(\alpha)$ vs. automation α , using a plausible simulated scenario.

C. Bias, Fairness, and Inclusion

Bias in artificial intelligence systems—especially in the context of agentic AI that makes decisions with little human involvement—has become a hot-button issue. Evidence of biased outcomes in large language models has permeated the media. The FBI Director has warned that biometrics authentication systems can “fail to be as accurate with Asian and African American faces as with Caucasian faces.” The global financial services firm Wells Fargo announced last year that it would no longer use facial recognition technology due to concerns about bias and the potential for wrongful arrest. Naturally, it would be unwise to dismiss these issues. Therefore, agentic AI solutions for customer service and claims processing must demonstrate inclusivity, fairness, and freedom from bias. Such evaluations need to occur in real-world multi-factor environments—not simply narrow testing contexts. Many of the areas in which AI is deployed can be sensitive. “The stakes are higher than ever,” as the U.S. Financial Services Regulatory Relief Act notes, with impacts on individuals and on society at large, particularly communities that have been historically underserved by the financial system. Consequently, established safeguards are imperative—both the external validation needed for a product, service, or agent and the internal monitoring required to keep it on an equitable path. Testing should be part of a product’s lifecycle and address risk across all relevant risk categories, from solutions specific to a financial institution to external risk reflected in industry models.

CONCLUSION

The proposed ideas furnish a logical understanding of the role of agentic decision-making in the claims processing and policy servicing initiatives. While these two activities have been defined separately, it is recognized that a large part of the flows could be integrated. In addition to the clear demar- cation of responsibilities within the second-level decisions, synergies during executions through modular agent design and orchestration have been briefly outlined. An important area of implementation- related research is explicated in the Governance, Trust, and Risk Management subsection, partic- ularly focusing on the auditability, appropriate allocation of accountability across insurer and customer, and guaranteeing appropriate safety margins. In the context of insurance claims and policy servicing, governance links back to answers to the questions of compliance and authority identified within the Ethical, Legal, and Regulatory Considerations section. Addressing the problem of operational risk across all areas, and therefore ensuring proper safety margins, is vital in build- ing customer trust. These trust-building requirements resonate with expectations of the insurance workforce and the broader society, and whenever agentic decision- making is applied in a context wherein external reinforcement factors hold, the requirements of auditability, explainability, and other aspects of governance assume lower priority. A reliable direction of future research concerns the integration of the flows for policy issuance, endorsements, and renewals with pricing adaptation and real-time underwriting, ensuring compliance along the way, especially in relation to fairness. To enable and encourage an equitable supervisory collaboration with insurance compa- nies and consequently promote the principles of trustworthy responsible AI, efforts should be geared towards Inclusive AI Design and all its sub-areas, with asset management firm- external technology providers looking for operational margin gains, especially through fraud detection, i.e. Insurance AI for whom insurance is not a core competency. Another ongoing research avenue focuses on the scalability of an agentic AI- powered ecosystem for claims processing, considering aspects such as governance, operational risk, and fairness.

A. Summary and Future Directions

The agentic AI foundations outlined here, and the asso- ciated governance framework proposed in Section 5, enable the next- generation insurance platforms to leverage agentic AI responsibly in claims processing and policy servicing. Agentic claims processing provides rapid response times, enhanced fraud detection and validation mechanisms, and process transparency. By reducing error rates, the risk posed to

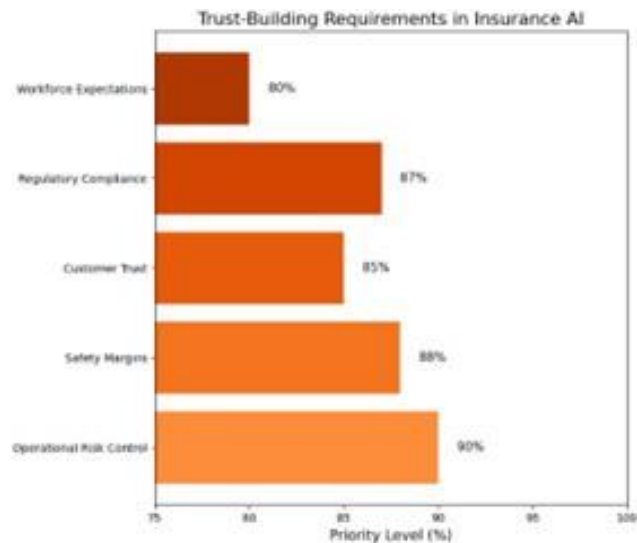


Fig. 13. Trust-Building Requirements in Insurance AI

The insurer is lowered. The opportunity to invest in high- stakes, low-probability risks is appreciated by customers. Careful balancing of redundant, explainable, and auditable processes will mitigate operational risk. Enabling automation in policy servicing enhances customer experience and satisfaction while optimally allocating labour. Research into the application of agentic AI for next- generation insurance platforms thus reveals considerable opportunities for enhanced customer experience and for lowering error costs, fraud detection costs, and other costs associated with process inefficiency in policy servicing and claims processing. Importantly, customer outcomes are enhanced by a principled application of agentic AI in these areas. The question of wider societal impact remains open and critical, and industry position suggests an important amplifying effect on other emerging technologies of

societal consequence. Addressing the systemic effects of increased automation in the workforce and across society requires continued engagement with the actuarial and insurance communities. A cross-cutting research agenda has been identified, encompassing gover- nance, risk and regulation considerations, technical challenges to safety and reliability, and the influence of agentic AI on other technologies with societal impact.

REFERENCES

1. Bhattacharya, Sukriti, German Castignani, Leandro Masello, and Barry Sheehan. "AI Revolution in Insurance: Bridging Research and Reality." *Frontiers in Artificial Intelligence*, vol. 8, 2025, article 1568266, <https://doi.org/10.3389/frai.2025.1568266>.

- Frontiers+1
2. *Beyond Automation: The 2025 Role of Agentic AI in Autonomous Data Engineering and Adaptive Enterprise Systems.* *American Online Journal of Science and Engineering (AOJSE)*, vol. 3, no. 3, 2025, ISSN 3067-1140, <https://aojse.com/index.php/aojse/article/view/18>.
3. Khayatbashi, S., V. Sjölin, A. Granåker, and A. Jalali. “AI-Enhanced Business Process Automation: A Case Study in the Insurance Domain Using Object-Centric Process Mining.” *arXiv*, 2025, arXiv:2504.17295.
4. Garapati, Ravi Shankar, and Suresh Babu Daram. “AI-Enabled Predictive Maintenance Framework for Connected Vehicles Using Cloud-Based Web Interfaces.” *Metallurgical and Materials Engineering*, 2025, pp. 75–88. <https://metall-mater-eng.com/index.php/home/article/view/1887>.
5. Indico Data. “Indico Data Launches Industry’s First Agentic Decisioning Platform Purpose-Built for Insurance.” *Indico Data Blog*, 4 June 2025.
6. McKinsey & Company. “The Future of AI in the Insurance Industry.” 15 July 2025.
7. Inala, R., and B. Somu. “Building Trustworthy Agentic AI Systems for Personalized Banking Experiences.” *Metallurgical and Materials Engineering*, 2025, pp. 1336–1360.
8. EIS Group. “A Complete Agentic AI Platform for Insurance — EIS OneSuite.” n.d.
9. Somu, B., and R. Inala. “Transforming Core Banking Infrastructure with Agentic AI: A New Paradigm for Autonomous Financial Services.” *Advances in Consumer Research*, vol. 2, no. 4, 2025.
10. Earnix. “Agentic AI Use Cases in the Insurance Industry.” *Earnix Blog*, 16 June 2025.
11. Meda, R. “Dynamic Territory Management and Account Segmentation Using Machine Learning: Strategies for Maximizing Sales Efficiency in a US Zonal Network.” *EKSPLORIUM–Buletin Pusat Teknologi Bahan Galian Nuklir*, vol. 46, no. 1, 2025, pp. 634–653.
12. Everest Group. “Agentic AI in Insurance: Transforming Risk, Relationships and Results.” 21 Nov. 2024.
13. Kummari, D. N., S. R. Challa, V. Pamisetty, S. Motamary, and R. Meda. “Unifying Temporal Reasoning and Agentic Machine Learning: A Framework for Proactive Fault Detection in Dynamic, Data-Intensive Environments.” *Metallurgical and Materials Engineering*, vol. 31, no. 4, 2025, pp. 552–568.
14. National Association of Insurance Commissioners (NAIC). “Artificial Intelligence (Insurance Topics).” 17 Jan. 2025.
15. McKinsey & Company. “The Future of AI for the Insurance Industry.” *McKinsey Insights*, 15 July 2025.
16. Sheelam, G. K. “Agentic AI in 6G: Revolutionizing Intelligent Wireless Systems through Advanced Semiconductor Technologies.” *Advances in Consumer Research*, 2025.
17. Shah, S. “Next Frontier: AI Driven Decision Making in Insurance.” *Fractal.ai*, 16 Jan. 2025; van Bekkum, M. “AI, Insurance, Discrimination and Unfair Differentiation.” *Journal of Risk and Insurance*, 2025, <https://doi.org/10.1080/17579961.2025.2469348>.
18. Yellanki, S. K., D. N. Kummari, G. K. Sheelam, S. Kannan, and C. Chakilam. “Synthetic Cognition Meets Data Deluge: Architecting Agentic AI Models for Self-Regulating Knowledge Graphs in Heterogeneous Data Warehousing.” *Metallurgical and Materials Engineering*, vol. 31, no. 4, 2025, pp. 569–586.
19. Debevoise & Plimpton LLP. “Europe’s Regulatory Approach to AI in the Insurance Industry.” 22 May 2025.
20. Zuiderveen Borgesius, F., M. van Bekkum, I. van Ooijen, G. Schaap, M. Harbers, and T. Timan. “Discrimination and AI in Insurance: What Do People Find Fair? Results from a Survey.” *arXiv*, 2025, arXiv:2501.12897.
21. Hill, G., J. Gong, T. Babeli, M. Mots’oehli, and J. G. Wanjiku. “LLMs and Agentic AI in Insurance Decision-Making: Opportunities and Challenges for Africa.” *arXiv*, 2025, <https://doi.org/10.48550/arXiv.2508.15110>.
22. Annapareddy, V. N., J. Singireddy, B. Preethish Nanan, and J. K. R. Burugulla. “Emotional Intelligence in Artificial Agents: Leveraging Deep Multimodal Big Data for Contextual Social Interaction and Adaptive Behavioral Modelling.” 14 Apr. 2025.
23. Luciano, E., M. Cattaneo, and R. Kenett. “Adversarial AI in Insurance: Pervasiveness and Resilience.” *arXiv*, 2023, arXiv:2301.07520.
24. Hogan Lovells. “Governance and Underwriting in the Age of AI: A Dual Challenge for Insurers.” 8 Oct. 2025.
25. Koppolu, H. K. R., R. S. Nisha, K. Anguraj, R. Chauhan, A. Muniraj, and G. Pushpalakshmi. “Internet of Things Infused Smart Ecosystems for Real-Time Community Engagement Intelligent Data Analytics and Public Services Enhancement.” *Proceedings of the International Conference on Sustainability Innovation in Computing and Engineering (ICSICE 2024)*, Atlantis Press, May 2025, pp. 1905–1917.
26. Stetler, N. “Reinsuring AI: Energy, Agriculture, Finance & Medicine as Precedents for Scalable Governance of Frontier Artificial Intelligence.” *arXiv*, 2025, <https://doi.org/10.48550/arXiv.2504.02127>.
27. Sheelam, G. K., H. K. R. Koppolu, and B. P. Nandan. “Agentic AI in 6G: Revolutionizing Intelligent Wireless Systems through Advanced Semiconductor Technologies.” *Advances in Consumer Research*, vol. 2, no. 4, 2025, pp. 46–60; Pandiri, L. “Exploring Cross-Sector Innovation in Intelligent Transport Systems, Digitally Enabled Housing Finance, and Tech-Driven Risk Solutions: A Multidisciplinary Approach to Sustainable Infrastructure, Urban Equity, and Financial Resilience.” *Proceedings of the 2025 2nd International Conference on Research Methodologies in Knowledge Management, Artificial Intelligence and Telecommunication Engineering (RMKMATE)*, 2025, pp. 1–12.
28. Meimandi, K. J., G. Arañguiz-Dias, G. R. Kim, L. Saadeddin, and M. J. Kochenderfer. “The Measurement Imbalance in Agentic AI Evaluation

- Undermines Industry Productivity Claims.” *arXiv*, 2025, <https://doi.org/10.48550/arXiv.2506.02064>.
29. “Infrastructure, Urban Equity, and Financial Resilience.” *Proceedings of the 2025 2nd International Conference on Research Methodologies in Knowledge Management, Artificial Intelligence and Telecommunication Engineering (RMKMATE)*, IEEE, 2025, pp. 1–12.
30. Everest Group. “Agentic AI in Insurance 2025.” 12 Sept. 2025.
31. Hill, G., J. Gong, T. Babeli, M. Mots’oehli, and J. G. Wanjiku. “LLMs and Agentic AI in Insurance Decision-Making: Opportunities and Challenges for Africa.” *arXiv*, 2025, <https://doi.org/10.48550/arXiv.2508.15110>.
32. Koppolu, H. K. R., A. L. Gadi, S. Motamary, A. Dodda, and S. R. Suura. “Dynamic Orchestration of Data Pipelines via Agentic AI: Adaptive Resource Allocation and Workflow Optimization in Cloud-Native Analytics Platforms.” *Metallurgical and Materials Engineering*, vol. 31, no. 4, 2025, pp. 625–637.
33. Mukherjee, A., and Hannah Hanwen Chang. “Agentic AI: Autonomy, Accountability, and the Algorithmic Society.” *arXiv*, 2025, <https://doi.org/10.48550/arXiv.2502.00289>.
34. Munich Re Automation Solutions. “New EU Act Regulates AI in Insurance: Rollout and Implications.” 30 Jan. 2025.