



Self-Governing AI Agents in Digital Banking

Niklas Andersson
Asst Professor

Abstract

Digital banking is fundamentally reliant on the Digital Automation of Banking Operations Workflows (DABOW). DABOW are governed by internal banking risk management policies, but they factor neither genetic tendencies, impersonation attempts nor socio-economic conditions through the resources traditionally provided by external actors. Expansion into end-to-end Digital Banking Workflow Automation (DBWA) allows the introduction of autonomous AI agents controlling DABOW and monitors by weaving external AI resources and Reasoning Services into the Banking Automated Notification Infrastructure for Handling External Requests (BANIHER) framework, ensuring natural language processing and robotic process automation ease. Conformity to global practices of Digital Banking extends the Banking Automated Notification Infrastructure for Handling External Requests – Transaction Processing (BANIHER-TP) design to encompass Customer Onboarding and Identity Verification (CO-IV) for secure user profile creation, Transaction Processing and Fraud Detection (TP-FD) for trusted transaction processing and Fraud Extraction (FE), all enabling service level performance.

Incorporating cyber intelligence services and infrastructure augments DBWA with Cognitive Data Management and Breach Prevention Monitoring to enhance response efficiency. The depth of technology reliance and tautology practiced by modern banking strengthens its appeal as a target for manipulation through model integration. Deployments as Division Chief Officers of Cyber Manipulation Networks, targeting the shaping of fake news using psy-ops, cooperate and coordinate the First Blockchain Design Acceptance, System Analysis Synchronization Simulation, and Predictive Analysis of Blockchain Networking Synchronization and Displacement Detection in Lab Design as visible. Pathology Closure Network status and Controls Enactment Network operational control reinforce adoption.

Keywords: Digital Banking, Banking Automation, AI Agents, Workflow Automation, Risk Management, Identity Verification, Fraud Detection, Fraud Extraction, NLP, RPA, Cybersecurity, Cyber Intelligence, Data Management, Breach Monitoring, Blockchain,



Predictive Analytics, Transaction Processing, Customer Onboarding, Digital Identity, Banking Infrastructure.

1. Introduction

Despite many digital transformations, banking operations still require human specialists who struggle to manage growing workloads efficiently. Market leaders have automated numerous repetitive processes using Automation® and Digital Workers. However, full-scale implementations remain rare, and human-in-the-loop operations remain essential, adding complexity and cost. Human-in-the-loop arrangements themselves have become emerging points of automation across industries. As a result, banks face a challenging paradox: while demand-surge transaction-based modules are price-takers, costs continue to rise in the workflow-rich closed-loop business modules where banks offer premium pricing. Regulatory expectations further complicate matters. Banks are expected to consider possible fraud on every transaction, including those involving micro-amounts in the simplest retail payments.

This dilemma calls for the development of fully autonomous AI agents capable of managing business workflows from end-to-end across integration points with utter digital ease. A conceptual framework is proposed to specify the roles and Responsibilities of autonomous AI agents in end-to-end banking Workflow Automation. Key requirements and essential building blocks are articulated. Specific customer onboarding and transaction-processing workflows serve as illustrative use cases. Synoptic definitions, critical capabilities, architectural patterns, levels of decision-making autonomy, and core risk-and-governance considerations are also discussed. Such agents would deploy a control strategy control CHARACTERISTICS which define the level of autonomy or the target-independent mode of operation □□□□□ less on an outer performance □ □/□ стий≥□ àëлмsда. The goal amid this potentially commercial blind spot is to facilitate fully digital banking services where the AI is customer-facing, omnipresent, autonomous in its own right, and continuously available.

1.1. Overview of the Study

The study aims to: (a) set the fundamental terminology and essential functions of an autonomous AI agent, (b) define an autonomous AI agent as an operational micro-service for digital banking, and, consequently, (c) establish the scope and implications of supporting end-to-end digital banking workflows with autonomous AI agents. The term workflow within this



context refers to the sequencing of diverse engineering tasks, each potentially addressed by distinct autonomous AI agents and possibly executed asynchronously. It is important to highlight that the scaling of digital banking operations and associated risks has reached a level that cannot be sustainably addressed by relying solely on human talent and expertise, and that thus technological processes should form the backbone of digital banking operations. The focus here is on the use of various types of AI-enabled technologies for technology-assisted or technology-driven solutions as defined earlier, and in particular on AI models that encompass the perception-planning-execution-presentation learning cycle.

The identified support processes are customer onboarding and identity verification, transaction processing and fraud detection. These processes have been richly documented and are critically relevant for banks and other financial institutions seeking to convert increased operational scale into improved profitability. For each of these processes, the support data requirements as well as the key decision criteria for all internal actors involved are specified. The intent is to portray these automation flows as operations supported by a multitude of AI-enabled components, where the sequencing of interventions constitutes the process workflow, with implicit or explicit control distributions across the involved actors and automated functions.

2. Background and Context

A wide set of automated services and workflows are available in modern Digital Banking offerings, delivered through the Digital Channels. However, much of the Bank's backend processing continues to be manual and prone to institutional bad practices such as excessive use of spreadsheets, low design standards, and gap risks during data transfer across different technology platforms. The first challenge is to automate those backend processes end-to-end through seamless interaction between the Digital Channels and the backend operations. Combining a single-source-of-truth Data Lake with workflow automation and alert functions, supported by Digital Channels, turns the Bank into a real-time Financial Services Hub that allows clients to use those financial products that best suit their needs. Automation must include Transaction Processing (Payments, Risk Management, Fraud Detection, and Reconciliation Functions) and the Customer Onboarding and Identity Verification Functions.

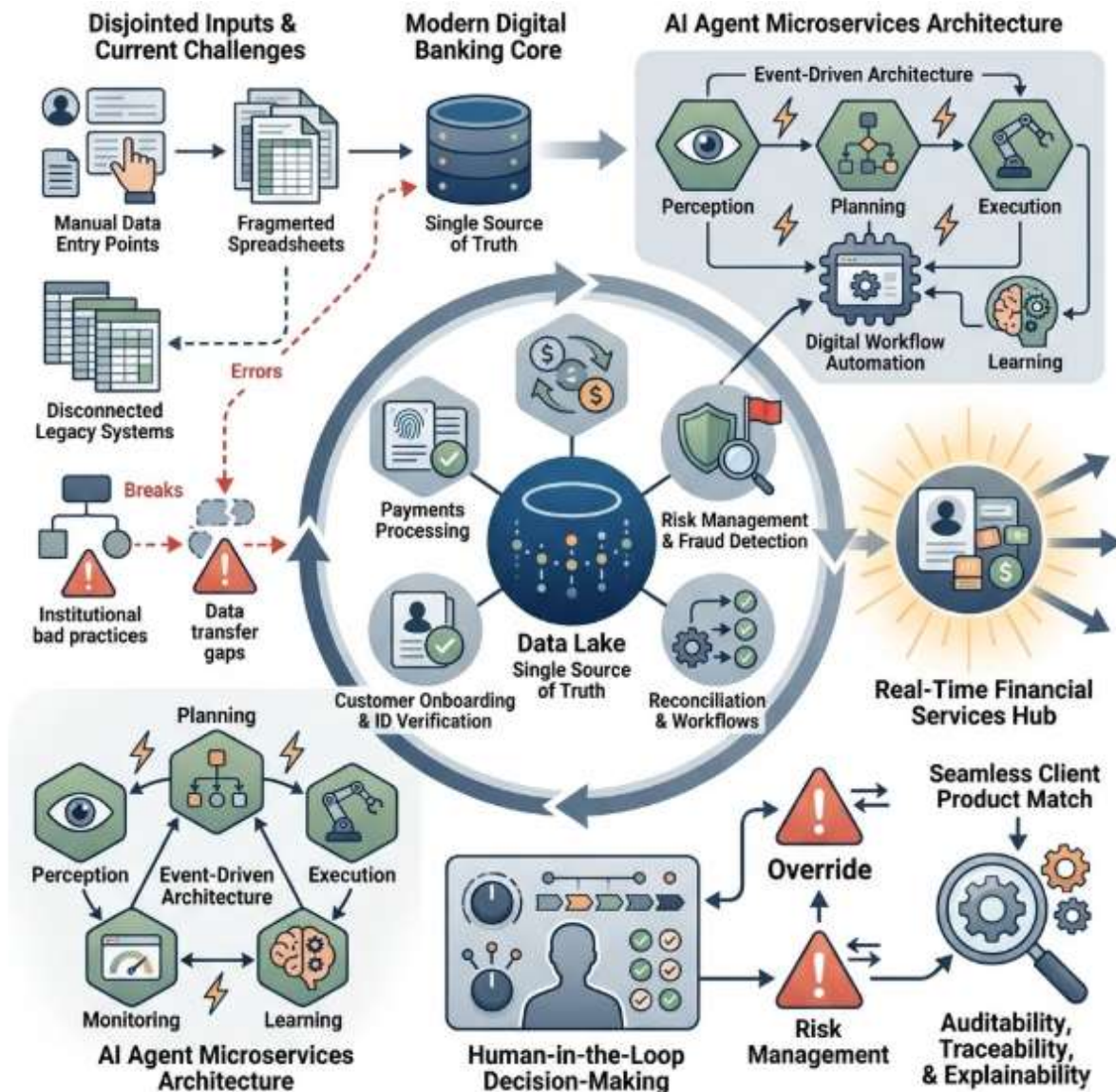
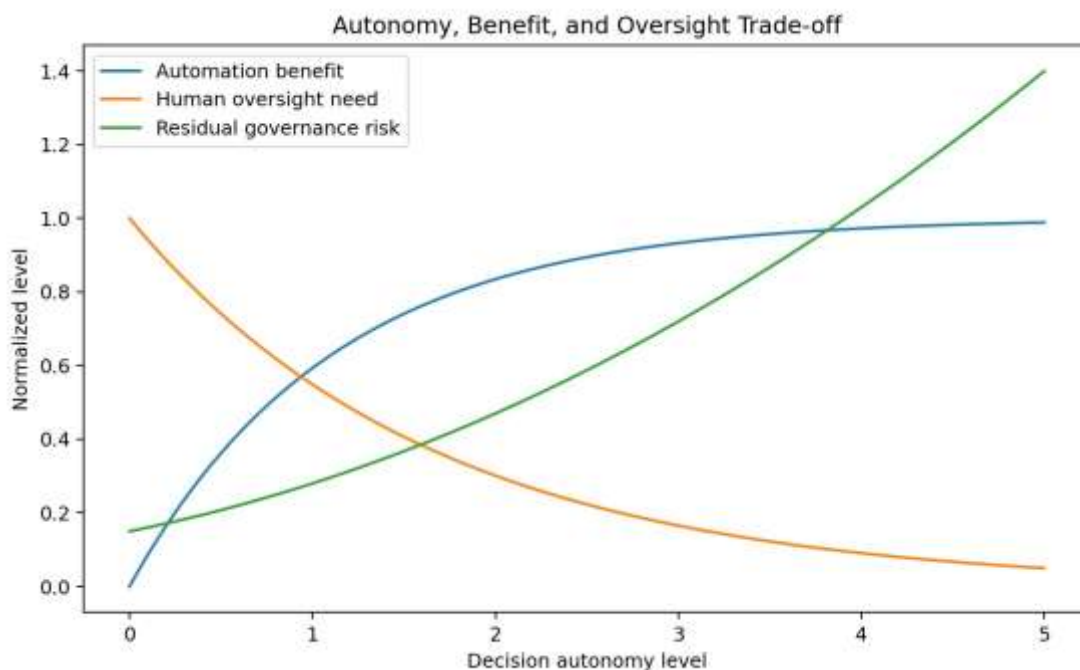


Fig 1: AI-Enabled End-to-End Banking Automation: An Event-Driven Microservices Architecture with Human-in-the-Loop Risk Controls

AI Agents — Software controlled by Artificial Intelligence (AI) with limited or no human control and capable of automatic planning, execution, and monitoring of simple or complex tasks — have emerged as a promising solution for End-to-End Automation of Digital Workflows. They must be designed to expose a full set of Core Capabilities: Perception, Planning, Execution, Monitoring, and Learning. A set of Microservices providing those Core Capabilities with an Event-Driven Architecture Pattern is recommended. Integration between Agents and/or automated Digital Workflows and the Bank’s Risk and Compliance Functions represent the main bottleneck for the adoption of autonomous AI Agents in banking. A Human-in-the-Loop Decision-Making Control Scheme for Risk Management and Fraud Detection Functions, based on Decision-Making Autonomy Levels, Decision-Making Thresholds, and Override Procedures, while ensuring Auditability, Traceability, and Explainability of all processes, represents the most effective approach to risk mitigation.



3. Methodology



The discussion on Autonomous AI Agents for end-to-end Digital Banking Workflow Automation can be divided into modular sections to facilitate comprehension. While the sections may be explored individually, the reader is best served by an in-depth reading of the entire work.

A set of five interrelated terms are essential for understanding the exploration of Autonomous AI Agents in the context of Digital Banking Workflow Automation. Autonomy denotes the ability to perform functions without human intervention. An Agent is a software application that performs tasks on behalf of a user. A Workflow is a series of operations performed in chronological order to achieve a specific business objective. An Integration Point is a connector that allows two independent systems to share a Data Stream. A Data Stream is a flow of data between two systems that share data that is useful and relevant, enabling each to enhance their respective services. Data Streams can be unidirectional or bidirectional, and data-sharing relationships can be either formal or informal.

The Banking Industry landscape is rapidly evolving, driven by the emergence of Digital Banks, progressive regulations, changing customer expectations, the accelerating relevance of AI technology in risk management and customer experience, and the need for operational efficiency and scalability. Yet consumer adoption of Digital Banks remains low, and the Consumer Financial Protection Bureau in the U.S. cautions that “the nation’s financial regulation should strive for the best compliance and create incentives for all forms of service provision and encourage healthy competition to provide the best guidance to consumers.” These insights highlight the need for thoughtful consideration and Investment in Compliance and risk management systems.

Equation 1: Customer onboarding composite risk score

Let

- x_1 = identity document confidence
- x_2 = proof-of-life confidence
- x_3 = AML/CFT compliance score
- x_4 = sanctions / watchlist clearance score



- x_5 = data completeness score

Let weights be w_1, w_2, w_3, w_4, w_5 , with

$$w_1 + w_2 + w_3 + w_4 + w_5 = 1$$

Then the **composite onboarding confidence** is

$$C = w_1x_1 + w_2x_2 + w_3x_3 + w_4x_4 + w_5x_5$$

Since banking decisions are usually expressed as **risk**, define

$$R_{\text{onboard}} = 1 - C$$

Step-by-step derivation

Start from the assumption that each onboarding sub-check contributes proportionally.

$$C = \sum_{i=1}^5 w_i x_i$$

Expanded:

$$C = w_1x_1 + w_2x_2 + w_3x_3 + w_4x_4 + w_5x_5$$

If confidence is high, risk should be low. So define risk as complement:

$$R_{\text{onboard}} = 1 - C$$

Substitute C :

$$R_{\text{onboard}} = 1 - (w_1x_1 + w_2x_2 + w_3x_3 + w_4x_4 + w_5x_5)$$

3.1. Fundamental Terminology and Essential Functions

The literature provides diverse perspectives on key concepts such as autonomy, agents, workflows, integration points, data streams, decision-making criteria, and control mechanisms. Beginning with the definition of agents, these are AI programs that act autonomously in response



to changes in the environment with minimum human intervention. Autonomy denotes the extent to which an agent requires human assistance, guidance, or supervision to function adequately. A human-in-the-loop model indicates that the agent depends, to a certain degree, on human oversight to confirm or complete its responses. Insight from risk management posits that an indifferent, clueless, or unaware fool might introduce incidents, deliberate malice might add to the incidents, and good individuals might prevent errors but cannot eliminate them altogether. The Nature of Customer Experience Web describes three layers of customer experience—transaction, contextual, and anticipatory—each becoming progressively more complex and challenging to manage.

Autonomous AI agents are capable of executing a complete end-to-end workflow without human intervention, are controllable through standard Operating Level Agreements (OLAs) and Service Level Agreements (SLAs), and adhere to local regulatory requirements pertaining to auditability, explainability, and traceability. Autonomy levels across agents engaged in end-to-end process automation follow the classic approach defined within the autonomous driving domain, where autonomy rises rapidly to a certain level but starts tapering off in the last mile. It is essential to know the level of decision-making autonomy for any action to ensure that the actions of the AI agent are in line with the intent of the individual and the individual retains control over crucial actions. Establishing decision-making autonomy thresholds is also vital for the definition of the human-in-the-loop mechanism and the escalation path for the function.

4. Objective of the Study

The preceding analysis explains that successful implementation requires well-defined business objectives, consideration of the autonomy level of the system and its components, a risk and benefit assessment at the appropriate level, and a technology stack and data architecture that fully support the agents.

The objective is to present a comprehensive conceptual model of AI-based agents and their application for automating end-to-end banking workflows. The direction, dimension, and risk profile of the banking decision-making process strongly influence the choice of automation and the level of supervision by people. Thus, understanding the decision-making process, the data available to estimate compliance, the integration points and the flow of data fundamental for these decisions, as well as the checks that have to be satisfied for an operation to be compliant, is a necessary step to develop and evaluate end-to-end banking transactions and processes. Such a



structured approach enables the development of automation rules that govern operational agents. These agents are autonomous systems capable of implementing a wide set of functions with little or no human involvement. They represent an evolution of traditional software bots, taking advantage of the recent advancement in AI.

Concept	Derived variable	Equation
Customer onboarding risk	R_{onboard}	$1 - \sum w_i x_i$
Fraud / anomaly detection	A	$\sqrt{\sum ((x_i - \mu_i) / \sigma_i)^2}$
Fraud probability	$(P(\text{fraud} A))$	
Human-in-the-loop threshold	$D(L, R)$	piecewise threshold rule
Autonomy	α	N_a / N_t
Reliability	R	$1 - N_f / N$
SLA compliance	S, S_c	threshold or normalized score
Throughput	λ	N / T
Average latency	$\bar{t}$	$(1/N) \sum t_i$
Drift trigger	Δ_M	(
Audit completeness	A_c	k / K
Expected loss	$E[L]$	$\sum p_i c_i$

5. Research Summary

The Design of Autonomous AI Agent in Banking is informed by the combined viewpoints of four intersecting areas including the Theory of AI Autonomy, AI Risk Management, Human-in-the-Loop AI, and Customer Experience in Banking. Central to all areas is the distinction between Automation and Autonomy, which clarifies the reasons for End-to-End Workflows, and ways of managing risk accordingly.

The Theory of AI Autonomy distinguishes between Low, Medium, and High Autonomy AI—where autonomy is defined by the degree to which a system can function independently



without the need for Human Intervention—to inform the evaluation of supervisory expectations. AI with Low Autonomy requires Human Intervention for any approval and cannot process transactions without Human Oversight during execution. Medium Autonomy refers to such models that can operate independently but the outcomes must be approved before proceeding to the next step. High Autonomy AI can perform end-to-end, automated processing but is constrained from executing specific actions based on parameters that automatically trigger a Governance condition.

5.1. Conceptual Foundations for Autonomous AI Agents in the Banking Sector

Four theories underpin the proposed architecture for autonomously intelligent agents in the banking sector. The first categorizes levels of autonomy in AI systems according to functions, supervision, and feedback loops. The second addresses the use of human-in-the-loop during the decision-making of apparent autonomous systems. The third relates AI autonomy and decision making to risk management and compliance. The fourth focuses on customer experience and trust.

The recent literature defines autonomy as the ability to act independently. An AI system is deemed fully autonomous when it perceives a situation, makes a decision, and executes a related action without human supervision or guidance. It lacks any feedback loop connecting it with a human being capable of assigning goals or giving instructions. Such fully autonomous AI systems, however, have not yet reached desired levels of accuracy. This has given rise to AI systems that are perceived as autonomous—sending users on a logical path from inputting data to outputting decisions, conducting the necessary processing but delegating decision making for specific cases to a human. The literature on almost– autonomous systems has established an architecture composed of five key components: perception, planning, execution, monitoring, and learning.

The success of these agents depends, among other things, on the proper definition and management of the decision-making autonomy territory—the operational territory in which the functioning is “self-governing to a specified and agreed level of autonomy.” Here, self-governing means “the capability of the system to determine actions by itself without supervision,” and the lower the autonomy level, the more supervision required from human decision makers and external governance agents, and the closer the system is to being a decision aid or business rule engine. This definition implies a wide and varying range of operational reaches and



responsibilities; indeed, with the appropriate safeguards, the same natural conditions that increase the trustworthiness and perceived fairness of an AI system support also a progressive expansion of the decision-making autonomy territory toward higher levels.

Autonomy level	Meaning	Mathematical control
Level 0	Human only	$T(L) = 0$
Level 1	AI recommends	very small $T(L)$
Level 2	AI executes low-risk cases	moderate $T(L)$
Level 3	AI handles most routine cases	higher $T(L)$
Level 4	AI autonomous with control gates	high $T(L)$, override active
Level 5	Full autonomy	theoretical upper bound

6. Conceptual Framework for Autonomous AI Agents in Banking

Synthesizing theories that address autonomy, the human-in-the-loop principle, customer risk and experience, technology needs, and suitable governance models yields a conceptual framework for autonomous agents. Perception, planning, execution, monitoring, and learning capabilities span the cognitive domains required of autonomous AI agents, while architectural patterns organize them into service layers for microservices, event-driven designs, and governance in accordance with established supervisory paradigms for AI and machine learning.

The technological transformation of banking heralds an exciting yet uncertain future. One key paradigm shift involves the use of autonomous AI agents, which combine machine learning, computer vision, natural language processing, user experience design, and both intelligent and traditional automation to minimize human intervention in the delivery of services. Autonomous agents can safeguard against the associated risks and enhance the customer experience if properly developed, deployed, governed, and controlled. Informed by contemporary theories relevant to banking, a conceptual foundation that represents these ideas much like a blueprint is offered here.

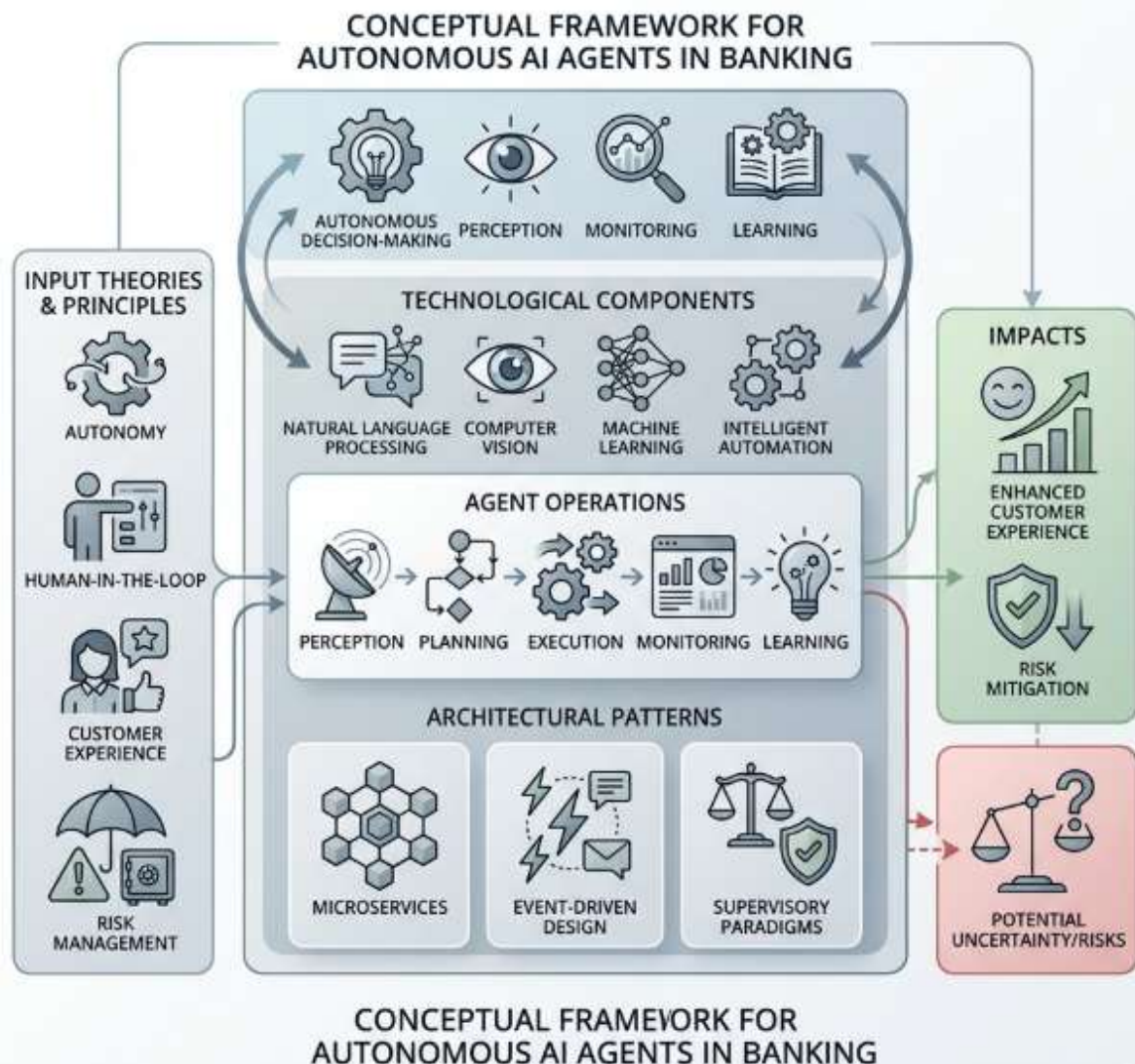


Fig 2: Towards a Cognitive Architecture for Governorable Autonomous Agents in Financial Services: A Foundational Framework

6.1. Definitions and Core Capabilities



High-performance autonomous AI Digital Agents, capable of undertaking the entire banking workflow from initiation to completion of service delivery, without any reliance on human intervention or oversight, cannot currently be implemented. Highly specific, focused agents, capable of functioning largely independently of other agents, can be developed, but these agents are not genuine Digital Agents, since they operate only within limited functional silos. End-to-end Digital Banking automation is an essential prerequisite for enabling the capability of Digital Agents, since their functionality is ultimately built on foundational high-performance AI models together with an effective orchestration mechanism. Autonomous agents lacking direct human oversight can nevertheless be deployed for limited-scope Digital Banking services, and the mechanism for building such agents is described.

Genuine Digital Agents require sophisticated human-like capabilities for vision, speech, language, general knowledge, reasoning, planning, and non-trivial task execution. Achieving the required levels of capability is an evolving process, and advanced implementations must therefore operate within an Autonomy Governance Framework that, at least for the time being, screens and oversees Executive Decisions. The Framework requires well-defined Operating Parameters, clearly identified Control Gates, failure response and escalation paths, and comprehensive audit, explainability, and model transparency mechanisms.

6.2. Architectural Patterns

Four fundamental architectural patterns will support deployment of autonomous AI agents across a banking environment.

Microservices Architecture for Breaking Workflows into Blocks: the need for rapid technology implementation, updates, fixes, and overall lifecycle management means that software developments must be planned and implemented in smaller manageable blocks, which can be fitted together to form the larger picture as a microservices-based architecture. Breaking down autonomous AI agents into microservices aids in ensuring that every segment can be developed, tested, and placed into production independently, which in turn reduces the time taken for new capabilities within the organization's broader software landscape.

Event-Driven Design for Systems Integration: banking is rarely a single event story, with most transactions being a consequence of multiple events. Therefore, a pull-based mechanism is more suitable than the conventional push-based approach so that complex autonomous agents



can act upon inputs from other environments and analysis engines rather than requiring other systems to inform them to be active. Using an event-driven design enables any change within the bank's ecosystem to be considered by the AI agents autonomously and, if needed, adjustments can be made to its working. Such an approach reduces the communication and transaction load throughout the bank's infrastructure.

Governance Layer with Control Interfaces for Watching Over Autonomy: while empowering systems to act autonomously can lead to improved efficiency and customer experience, such independence comes with an inherent risk. To mitigate this risk, the appropriate technology and process governance measures need to be in place to ensure the agents operate within accepted boundaries. It is crucial that the AI agents have clearly defined operating levels, thresholds, and human communities that can intervene when necessary, along with sufficient logging and audit trails.

Equation 2: Fraud anomaly score

Let a transaction vector be

$$\mathbf{x} = [x_1, x_2, \dots, x_n]$$

Let the expected normal profile be

$$\boldsymbol{\mu} = [\mu_1, \mu_2, \dots, \mu_n]$$

Let variability be represented by standard deviations σ_i .

Define standardized deviation for feature i :

$$z_i = \frac{x_i - \mu_i}{\sigma_i}$$

Then an aggregate anomaly score is

$$A = \sqrt{\sum_{i=1}^n z_i^2}$$

Step-by-step derivation

For each feature:

$$z_1 = \frac{x_1 - \mu_1}{\sigma_1}, \quad z_2 = \frac{x_2 - \mu_2}{\sigma_2}, \quad \dots$$

Stacking them:

$$A^2 = z_1^2 + z_2^2 + \dots + z_n^2$$

So:

$$A = \sqrt{z_1^2 + z_2^2 + \dots + z_n^2}$$

Substitute back:

$$A = \sqrt{\left(\frac{x_1 - \mu_1}{\sigma_1}\right)^2 + \left(\frac{x_2 - \mu_2}{\sigma_2}\right)^2 + \dots + \left(\frac{x_n - \mu_n}{\sigma_n}\right)^2}$$

7. End-to-End Banking Workflows: Scope and Implications

Despite extensive indications of banking's transformation among various stakeholders, the processes of automation remain inconsistent with much of the completed work addressed at point automation rather than at process flow. Yet only through the automation of holistic FinTech-level processes can true efficiency be achieved, along with meeting customer mission-critical requests. The recent rise in investment to build AIs with higher levels of autonomy may develop into a solution to the problem of point automation, enabling the end-to-end automation of digital banking processes. Examples that return fulfilment processes subjected to regulation and require human sign-off at various stages are customer on-boarding for identity verification/money laundering prevention and transaction fulfilment for fraud detection.

Through a mapping of data sources and flow, along with the decision trees for fulfilment processes, sufficient data and verification steps are indicated to support autonomous "AI" level fulfilments, breaking them down into sub-processes that use AI for detection but require greater fail-safe/protection and regulatory oversight—particularly for event monitoring, evaluation,



approval/disapproval, and auditability of controls, risks, and customer experience—forming part of a human-in-the-loop paradigm. Customer on-boarding receives particular attention, as completion of the full checklist is usually required for a successful fulfilment, with trust concerns particularly prevalent in this area.

7.1. Customer Onboarding and Identity Verification

Autonomous AI Agents for End-to-End Digital Banking Workflow Automation

Gaps in digital banking workflow design and autonomous agent capabilities have hampered banks' ability to realise the full potential of automation. Agents, defined as AI systems that autonomously perform tasks, can only provide automation for workflows if they possess the required data and expertise. Banks can realise end-to-end digital automation by deploying autonomous agents capable of onboarding customers and processing transactions without human intervention. Plugging the gaps that hinder banks from achieving this outcome requires an understanding of the data that the agents use and of the tasks involved.

Customer onboarding and identity verification—a sub-process of onboarding—provide a suitable starting point for assessing technical feasibility and design implications. Customer-onboarding data lineage can be mapped to the upstream sources of proof-of-life and identification information, the information-requirements within the applicant risk-scoring process, and the data required for regulatory compliance checks. Completing a data-mapping exercise for identity verification can help surface agent-decision criteria and consolidate the various proofs into a single integrated stage of the risk-engine process.

Autonomous AI Agents for End-to-End Digital Banking Workflow Automation

Gaps in digital banking workflow design and autonomous agent capabilities have hampered banks' ability to realise the full potential of automation. Agents, defined as AI systems that autonomously perform tasks, can only provide automation for workflows if they possess the required data and expertise. Banks can realise end-to-end digital automation by deploying autonomous agents capable of onboarding customers and processing transactions without human intervention. Plugging the gaps that hinder banks from achieving this outcome requires an understanding of the data that the agents use and of the tasks involved.

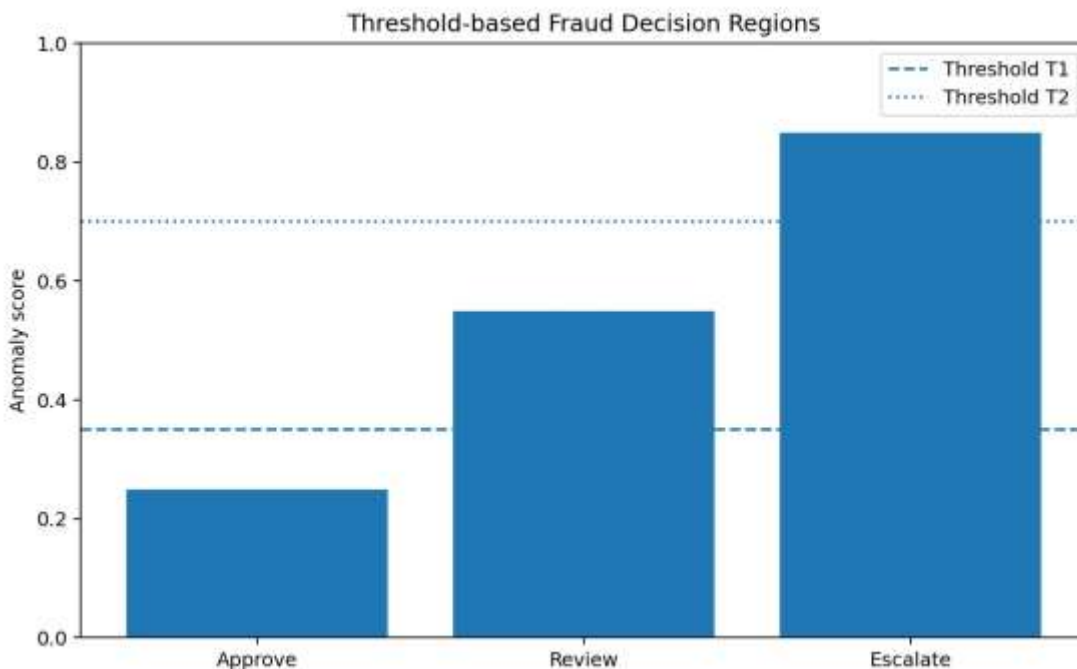


Customer onboarding and identity verification—a sub-process of onboarding—provide a suitable starting point for assessing technical feasibility and design implications. Customer-onboarding data lineage can be mapped to the upstream sources of proof-of-life and identification information, the information-requirements within the applicant risk-scoring process, and the data required for regulatory compliance checks. Completing a data-mapping exercise for identity verification can help surface agent-decision criteria and consolidate the various proofs into a single integrated stage of the risk-engine process.

7.2. Transactions Processing and Fraud Detection

Data architecture and data quality concerns also underpin the second end-to-end banking workflow being explored: transaction processing with fraud detection. Data must flow through a pipeline of processing stages, much like a manufacturing assembly line. In this case, the final product is a set of confirmation messages sent to parties involved in the transaction. Whenever something goes wrong, a separate error-handling process must step in. Several performance metrics must be tracked, manually scaled to demand, and kept within defined service-level objectives. A clear definition of expected bank behaviour—as opposed to customer behaviour—forms the foundation for the anomaly-detection models, which are trained to detect patterns that deviate from the norm.

A precisely mapped transaction-processing pipeline makes it possible to add a degree of automation and oversight to a bank's fraud-detection process. Anomaly-detection models watch over transaction-processing operations and flag any output that strays from the norm. Depending on the severity of the anomaly detected, separate groups of trained administrators can either approve flagged transactions directly or triple-check them before giving the green light. In either case, anomalies that pass through are automatically captured and retained for eventual reconciliation against detection model predictions. In other words, validation of bank confirmations becomes a second use case for predictions generated by anomaly-detection models.



8. Autonomy Levels, Governance, and Control Mechanisms

Decision-making autonomy of autonomous AI agents can be classified into five levels, which determine the extent of human oversight needed in the agent's operation. The decision-autonomy level of each workflow control point is specified. At points where the allowable autonomy is not sufficient to handle an exception, a process is defined for human override or escalation. The levels are summarized in Tab. 8.1.

Decision-Autonomy Levels

The use of machine learning to support autonomous decision making raises questions of auditability, explainability, and traceability of the decisions being made. Auditability ensures that the decisions made by models can be inspected and that those decisions adhere to reasonable rules. Explainability ensures that when the models make a decision, it can be explained why that decision was taken. Traceability ensures that it is possible to see how the decisions being made



by the models can lead to bias toward certain groups.]” When managing different classes of risk, auditability, explainability, and traceability can be handled in different but related ways.

Equation 3: Probability of fraud from anomaly score

Banks often need probability, not just raw anomaly score.

Use logistic conversion:

$$P(\text{fraud} | A) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 A)}}$$

Step-by-step derivation

Let the linear score be

$$y = \beta_0 + \beta_1 A$$

To convert any real number y to a value between 0 and 1, apply sigmoid:

$$\sigma(y) = \frac{1}{1 + e^{-y}}$$

Substitute y :

$$P(\text{fraud} | A) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 A)}}$$

8.1. Decision-Making Autonomy and Human-in-the-Loop

Decision-making autonomy encompasses different capability levels across seven stages, active agent interaction with the outside world, explicit separation of technical and substantive decision realms, human confirmation at different process phases, and the electronic circulation of non-critical decisions or processes. Adverse events may trigger management alerts or control intervention. The human-in-the-loop mechanism operates as an extra safety net; human review and confirmation remains a regulatory requirement for specific business domains or sensitive activities. In cases where human supervision falls short or diminishes user confidence, regulators



may opt to provision increased oversight during pre-production control cycles. Supervisors or users may establish consistency targets, requiring decisions made outside control thresholds for certain topics or areas to be spot-validated or unit-validated.

Approval and validation processes reflect the potential human-in-the-loop nature of the service. The automated processing pipeline outlines when workflow steps require human supervision. Approvers' recommendations circulate again to the originator after completing the additional steps. Supervisors may then define an escalating approval behavior: subsequent approval or rejection by a more senior manager. Such a process minimizes the risk of overloading senior management capacity in organizations with extensive growth during an economic boom cycle. Process auditability enables the tracing of all request paths. Improvement provisioning leverages the knowledge acquired through user confirmation. Updating and versioning structures enables constant training of the Automated AI Governance model.

8.2. Auditability, Explainability, and Traceability

Comprehensive auditability and traceability build user trust, facilitate external scrutiny, and help identify weaknesses. Each decision made by autonomous agents must be logged and preserved in compliance with applicable laws and regulations, creating an audit trail. Subject to the degree of automation and the associated residual risk, archives should record a comprehensive overview of data ingestion and key modules and systems involved in the decision. Audit evidence must be of sufficient granularity and detail to allow understanding by a suitably qualified person, including model, data, and result provenance.

Full transparency of AI models might not be possible owing to trade secret protections, national security, or other overriding interests. Regulations sometimes require explanations that are human-understandable. Popular techniques for interpretable model-agnostic explainability include LIME, SHAP, and counterfactual explanations; these can be complemented by increased user engagement and education. Where full transparency cannot be provided, financial institutions must establish appropriate justification.

Traceability relates to the ability to follow the lifecycle of a transaction from inception to completion, including its treatment by various financial institution systems and the changes made at each stage. The human-in-the-loop strategy contributes to transparency for users: requests for sensitive transactions to be authorized by a responsible individual or committee may be captured



in a system external to the AI-driven process. Such records can complement logs of the autonomous agents processing undertaken.

9. Technology Stack and Integration Strategies

The technology stack of autonomous AI agents in digital banking must consider banking regulation and supervisory practice. Banking supervision requires adherence to well-established regulatory paradigms, including risk-based supervisory expectations. Existing laws, regulations, and frameworks for governance, risk management, ethical use, AI explainability, bias mitigation, and data privacy must therefore be addressed in its construction. Robust and effective implementation further depends on sound pragmatic data architecture with verifiable data quality, a rigorous and reliable model lifecycle management capability for third-party and in-house models, and well-defined procedures for incident response and other operational metrics.

A well-defined banking data architecture enables banking supervisory authorities to perform effective supervision. Precise and consistent data models are therefore essential. Data cleaning processes should follow established data governance frameworks and must provide the necessary assurances on data privacy within respective jurisdictions. Access to sufficient quality and quantity of training data for deployment of new autonomous agents should be clearly defined. Transparent model lifecycle management for the development, operationalization, and retirement of all in-house and third-party models is critical. Robust logging, audit trails, explanation of model decisions and predictions, and traceability of supervisory information back to model inputs are essential for effective supervision.

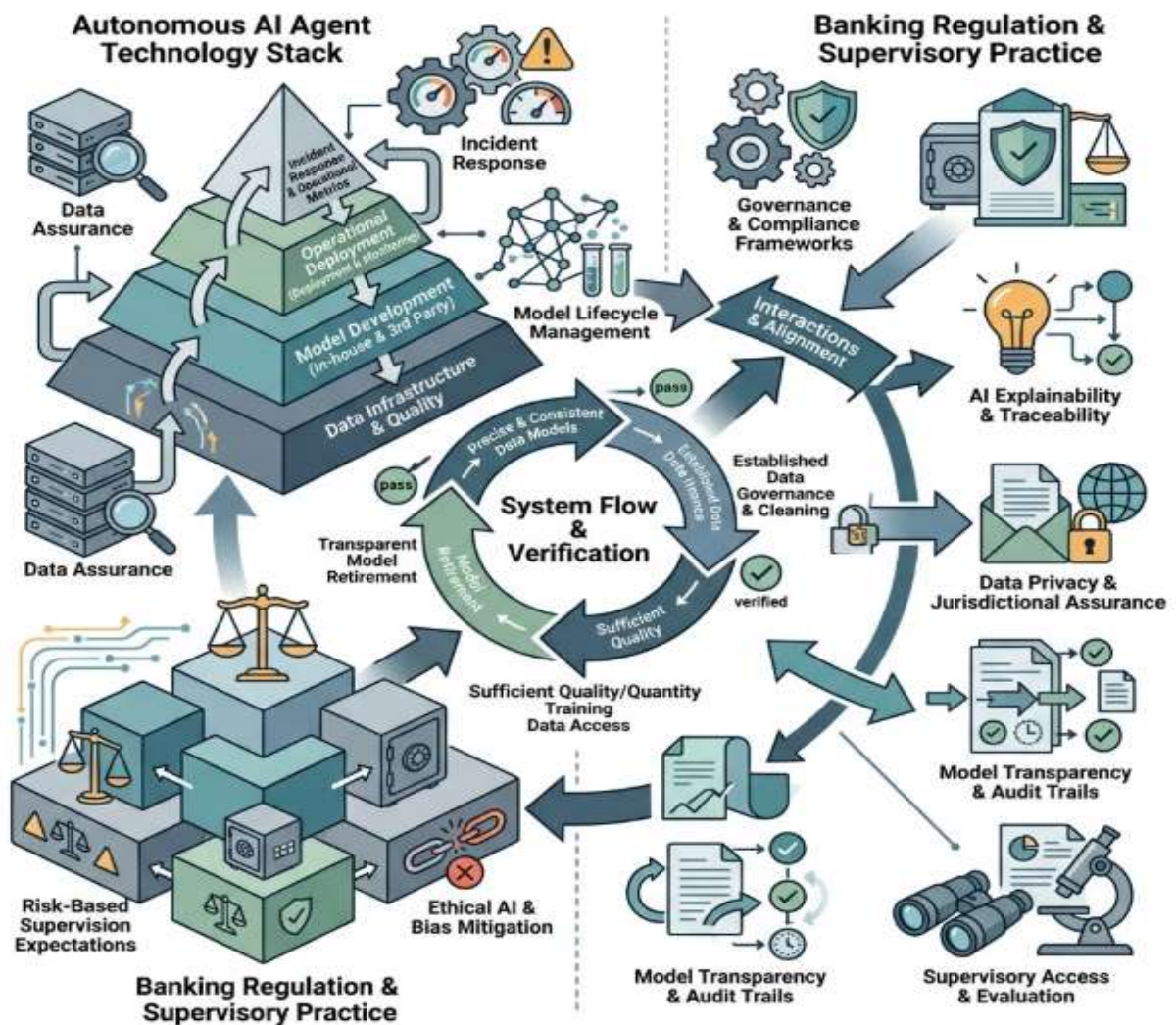


Fig 3: Governing Autonomous AI Agents in Digital Banking: Aligning Technology Stacks with Regulatory Paradigms

9.1. Data Architecture and Data Quality



An integrated technological infrastructure is essential for autonomous banking agents. An agent capable of end-to-end customer onboarding at a bank typically requires access to data from multiple systems, knowledge of data governance regulations, and fulfilled data quality requirements throughout the information supply chain that steers the life cycle of a transaction. Data governance is the procedure that ensures data integrity, compliance, and auditing. It entails defining policies, assigning roles, establishing rules, and implementing control processes supporting the proper management of data flows throughout an organization. Such control is essential to establish appropriate data lineage.

A superior quality of the data that flows through an information supply chain considerably reduces the chances of the models that are built on top of such data producing poor predictions. The data cleansing processes help remove spurious information. The disclosure of non-vital personal information must be restricted by means of proper anonymization techniques whenever the data are used in training models. Another aspect that impacts data quality is the implicit bias that is transmitted with the data. Such bias must be managed with dedicated fairness and de-biasing techniques both to prevent future unfair decisions and to increase customer trust in AI-based tools and their predictions.

9.2. Model Lifecycle Management

Management of the model lifecycle is essential to maintain optimal performance. Models must be re-trained when their data distribution no longer corresponds to the training data distribution (input shift) or when changes occur in the distribution of the target variable (target shift). One method to adapt to input shift is to monitor key model performance metrics (e.g., accuracy for classification, RMSE for regression) and re-train the model if the set thresholds are exceeded. Models may also be re-trained as an ongoing process—for example, at a definite time interval, for instance, once a week or month, an equivalent period in the training data being ensured. In the event of target shift, models must be retrained whenever sufficient labeled historical records are made available. In various domains (e.g., computer vision), model improvements often come from training with larger model architectures on larger datasets that are laboriously annotated. In this scenario, a strategy can be employed whereby the model version is changed, for example, to a larger architecture, each time a dataset of sufficient sizes is obtained.



Once a new version of a model is prepared for deployment, its predictive ability on the corresponding test dataset must be monitored to ascertain that it exceeds those of all currently deployed versions. Following deployment, the predictive ability of the new version can also be checked against that of the previous version, permitting easy roll-backs of its deployment if necessary. Finally, models that have been recently trained remain in storage until sufficient data for a newer model version are available, the actual age of the model version in production being factored into the choice. By these means, the use of an inferior model version can be minimized. The model-lifecycle considerations outlined above are conceptually similar to the continuous-delivery pipeline in software-development, where stages of the pipeline mimic the steps taken during product development.

Equation 4: Human-in-the-loop decision function

Let

- L = autonomy level
- R = risk score
- $T(L)$ = maximum risk that level L may handle automatically

Then the decision function is

$$D(L, R) = \begin{cases} \text{Auto-execute,} & R \leq T(L) \\ \text{Escalate to human,} & R > T(L) \end{cases}$$

Step-by-step derivation

At each autonomy level, there is an allowed operating territory.

If system risk is inside that territory:

$$R \leq T(L)$$

then autonomous execution is allowed.

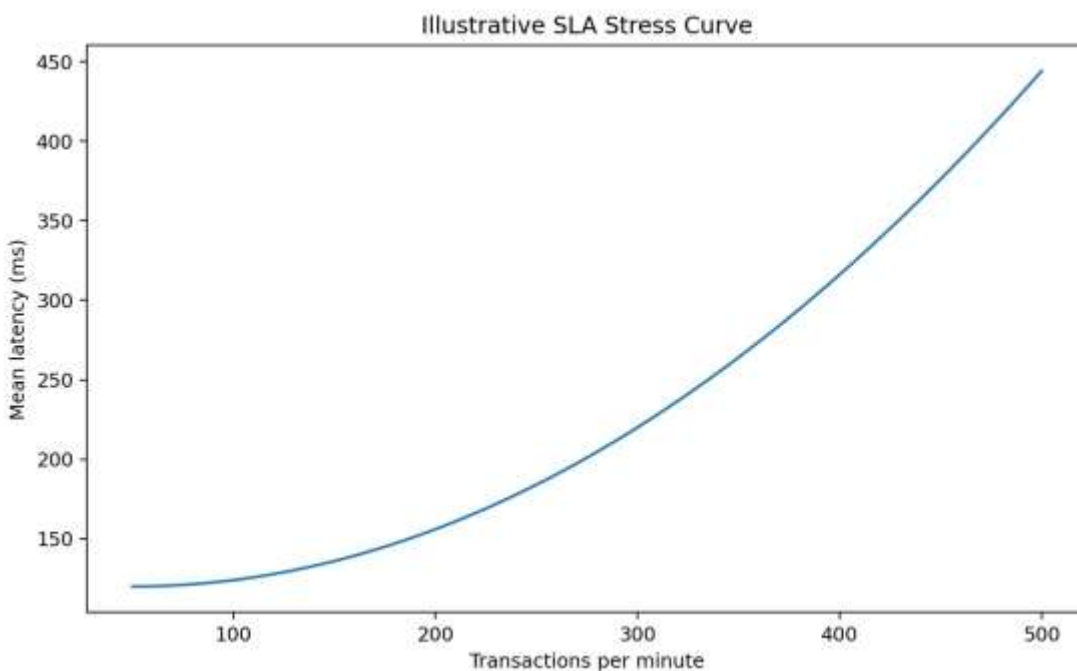
If risk exceeds it:

$$R > T(L)$$

then human intervention is required.

So piecewise:

$$D(L, R) = \begin{cases} \text{Auto-execute,} & R \leq T(L) \\ \text{Escalate,} & R > T(L) \end{cases}$$



10. Risk, Compliance, and Ethical Considerations

Five distinct aspects of Regulatory Paradigms and Supervisory Expectations affect the Evaluation Stage for Risk, Compliance, and Ethical Considerations. First, many jurisdictions lack comprehensive legislation or regulatory guidance specifically addressing AI systems. Consequently, AI-based solutions operating in such regions remain largely unregulated from a technology perspective and may be deployed without significant oversight. Second, areas of the



AI implementation stack that are either fully or partially regulated (such as payment systems, securities trading, and financial services) often present underlying risks that are assessed, monitored, and mitigated without reference to the AI technology enabling their operation. In such cases, the AI may be treated as any other technology underpinning a significant piece of regulated financial infrastructure. Third, the technology is often considered too immature for use in safety-critical applications, such as self-driving cars, advanced autopilots in commercial aircraft, and the management of complex nuclear power plants; hence, the deployment of AI in these areas is either prohibited or substantially delayed.

Fourth, expectations arise from the concept of Responsible AI, which proposes that AI technologies must operate in a fair, auditable, accountable, transparent, secure, trustworthy, privacy- and data-protective, and environmentally-friendly manner. Fifth, trust in AI systems remains critical from both customer and market perspectives, and emerging AI bias along any dimension—whether direct or indirect—may severely damage user confidence. Such risk is compounded whenever an automated AI decision is either not made with the consent of the affected customer or is unaccountable in any way.

10.1. Regulatory Paradigms and Supervisory Expectations

Data and services in the digital banking landscape are heavily regulated for several reasons, including money laundering prevention, consumer protection, and criminal activity deterrence. These regulations directly affect where and how data may be stored and processed, both technically and contractually. Given the amount of personally identifiable information in entering and operating in the banking system, digital banks also devote considerable resources to ensure that customer onboarding is as seamless as possible. This is enhanced by wizard-style forms exposed via web or mobile channels. Regulatory pressure leads to consideration of all information required when the customer is first entered into the system; filling out empty “nice to have” fields takes minimal effort.

The banking sector has embarked on a journey to offer frictionless customer onboarding while maintaining compliance with anti-money laundering (AML) legislation and associated requirements imposed by supervising agencies for combating money-laundering (CFT) purposes and fraud detection. A novel banking solution is PPP – People-Pay-Process – geared toward 100% digital banks. It allows frictionless onboarding without compromising speed, reliability, or regulatory compliance. Unlike traditional banking systems, the concept ensures that only data



absolutely necessary for the first transaction is collected while optimizing all data into the process pipeline for future use in CFT and AML activities.

10.2. Ethical Implications and Customer Trust

Trust is paramount to customer acceptance and adoption of banking services. Bias in automated processes can lead to customer disenchantment and subsequent loss of business. Such situations can arise, for instance, where risk classification is based upon inadequate perception training data or where machine-learning models are built upon training data whose development includes bias. Autonomy of AI-based digital banking services at the recommendation-to-satisfaction layer can create situations where the Trade-Off Theory applies. The SLA between Customer and Bank Service addresses the bank Customer's need for perceived Trust, a virtual Good, and Trade-Off POV therefore becomes a key element on the G2C side.

Trust in any service is predicated by the service increasingly meeting expectations. Autonomous AI Agent Services would therefore require a system being in place such that bias in training, infrequent decision-model re-training, and active personnel involvement and monitoring is organized and finalized to sufficiently meet expectations of the Customer quickly. A Service Level Agreement System between the Bank and Customer is therefore mandatory to place the necessary controls (Active-Human monitors responsible for actually monitoring AI-based Service decisions). However, how closely such decision outcomes need to be monitored is only practical for specific cases where banks have shown sentiment of extreme caution. Balancing the concept with Trade-Off would create Banking Services that can gain the Customer's satisfaction.

11. Evaluation, Metrics, and Validation

Metrics must be defined to evaluate whether the autonomous agents' objectives are achieved. Performance metrics consider the execution of specific tasks or workflows, while operational metrics are concerned with the availability and support for those tasks. Operational metrics include all service-level agreements (SLAs) defining the service levels, metrics assuring internal quality and reliability, and those measuring the output of functionally dependent systems. Metrics for the autonomous agents should address external performance targets, including accuracy of decision-making, processing latencies, throughput, and reliability.

Two lists of performance measures can be established for a bank's autonomous agents: primary metrics that assess their primary functions and the quality of execution, and secondary



metrics developed for human-in-the-loop, risk management, and customer experience. Auditability, explainability, and traceability form the basis of those measures that address the requirements of internal and external scrutiny for risk management and compliance. Decision-making autonomy establishes granular coverage for the authorizations that move from automated to supervisor-controlled and, subsequently, to human-only delegation. For bank management, the reliability of the agents represents the most important outcome of their use; maintaining low support effort is the top priority for the responsible business unit.

11.1. Performance Metrics for Autonomous Agents

Performance metrics for autonomous agents operationalize the qualitative design goals into quantitative targets, guiding the success of development efforts. Direct action autonomy is the ability to make decisions without human intervention; decision-making empowerment reflects the extent of human-endorsed actions. Ideally, actions are unrepeated and timely. Key process characteristics include execution speed, transaction velocities, and detection delays; key outcome characteristics include accuracy and reliability of decision-making, success and failure rates of task accomplishment, and incidence of adverse consequences.

Autonomous agents represent a revolution in the banking archetype, relieving banks of executing mundane tasks while enabling clients to conduct any kind of transactions—all at a considerably reduced cost and improved experience. Such agents enter the market heavily based on recommendation engines, fraud detection, internet searches, or chat support applications. The supervisory stance remains vigilant: afforded independence, considered automatically in the absence of clear deferral to human judgment, such agents may dictate or influence decisions with notable consequences even extending to compliance.

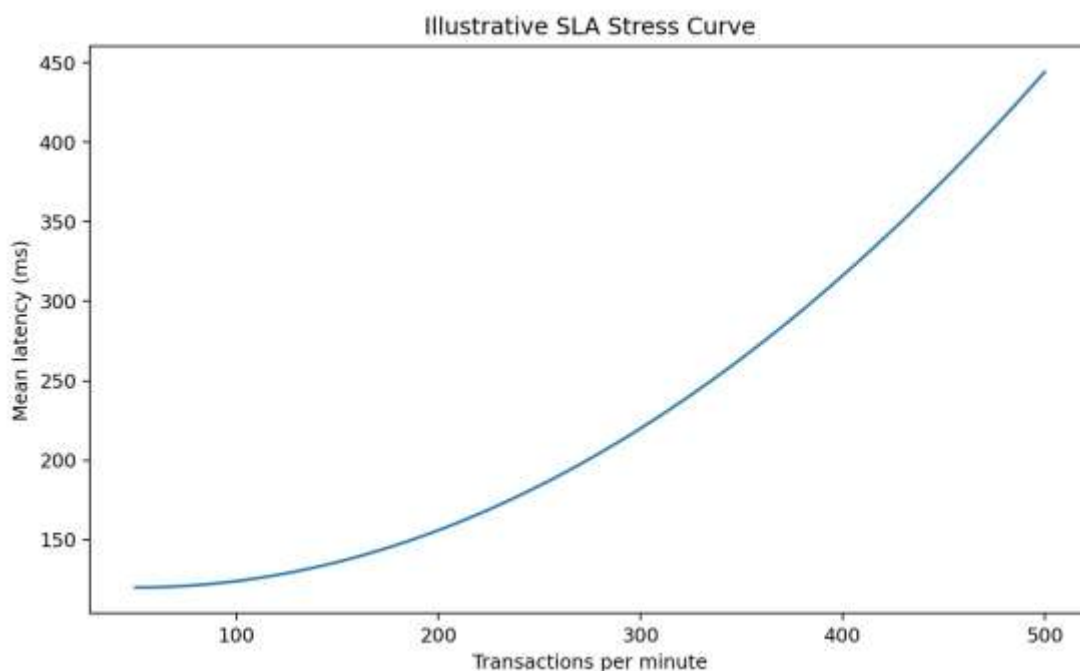
11.2. Operational Metrics and Service Levels

Many operational attributes are key factors in decision making, including resiliency, availability, minimum time to recover from incidents, incidents frequency and type, error rates, sensitivity, responsibility and monitoring. Such attributes often require service-level agreements (SLAs), which define a contractual relationship between a service provider and its customers. SLAs help assess the performance of service providers and set up penalties for non-compliance.

Stricter SLAs may depend on who is consuming the service. For example, an anomaly detected by the fraud detection system that has a low probability of being false (very low risk of



being false) could be automatically approved, but it might still require further verification if it is flagged by the compliance-monitoring system. Other SLAs are based on the cascading effect of managing multiple services together. If a delay occurs in one service, it could affect other dependent services downstream. For example, if a service that generates beauty reports for a business account is delayed, the reports' external distribution should be delayed as well.



Equation 5: Autonomy score

Let

- N_a = number of actions completed automatically
- N_t = total actions

Then a natural autonomy ratio is



$$\alpha = \frac{N_a}{N_t}$$

Step-by-step derivation

If every action is automatic, autonomy should be 1.
If no action is automatic, autonomy should be 0.

That directly implies:

$$\alpha = \frac{\text{automatic actions}}{\text{all actions}}$$

So:

$$\alpha = \frac{N_a}{N_t}$$

And human-dependency ratio is

$$H = 1 - \alpha$$

12. Implementation Roadmap and Case Studies

Staged deployment strategies mitigate risks for critical supervision-focused workflows, and tandem rollout and evaluation enable other workflows' validation and scaling. A customer onboarding agent validates overtime policy adherence and helps build customer trust.

A risk-based approach to implementing autonomous agents affords adequate scrutiny without compromising deployment velocity. Regulatory approval or pilot success instills confidence at subsequent stages; pilots reap reward or rollback principles if ruled dangerous, with gradual or abrupt downsizing depending on supervisory tolerance for automatic retreat from automation.

Changing authorization and regulatory engagement thresholds fosters user and regulator confidence. Pilot agents fulfill an implicit human-in-the-loop paradigm, with their recommendations deemed non-binding even if regulatory consultation occurs. Customers receive



guarantees of false positive-free service as a condition of successful pilot completion. Experience-building improvements in decision speed and data availability chronicle case closures, encouraging self-indulgent client onboarding while promoting investigators' learning internalization.

Risk supervision is notoriously laborious, and the exposure thus earned fosters user-focused product merits rather than supplier-centric focus. Agentization proves warranted under principle of least effort: deploying a dedicated agent eases effort on supervisors by processing proposals that present little likelihood of refusal within supplemented capabilities. The risk process inherently demands scrutiny; deploying a pilot agent refines its operation without accelerating operational pace.

In practice, the filtering-enhanced process operates over legacy systems, with naturally integrated enriched data streams preparing contextual bandwidth for scaling to supervisory requirements. Promoting risk scoring and loyalty assessment processes to supervisory data streams enhances processing speed via regularized data formats and noise-removal interventions; doing so frees bandwidth for experts to manually tackle historical patterns or renew classifier scope.

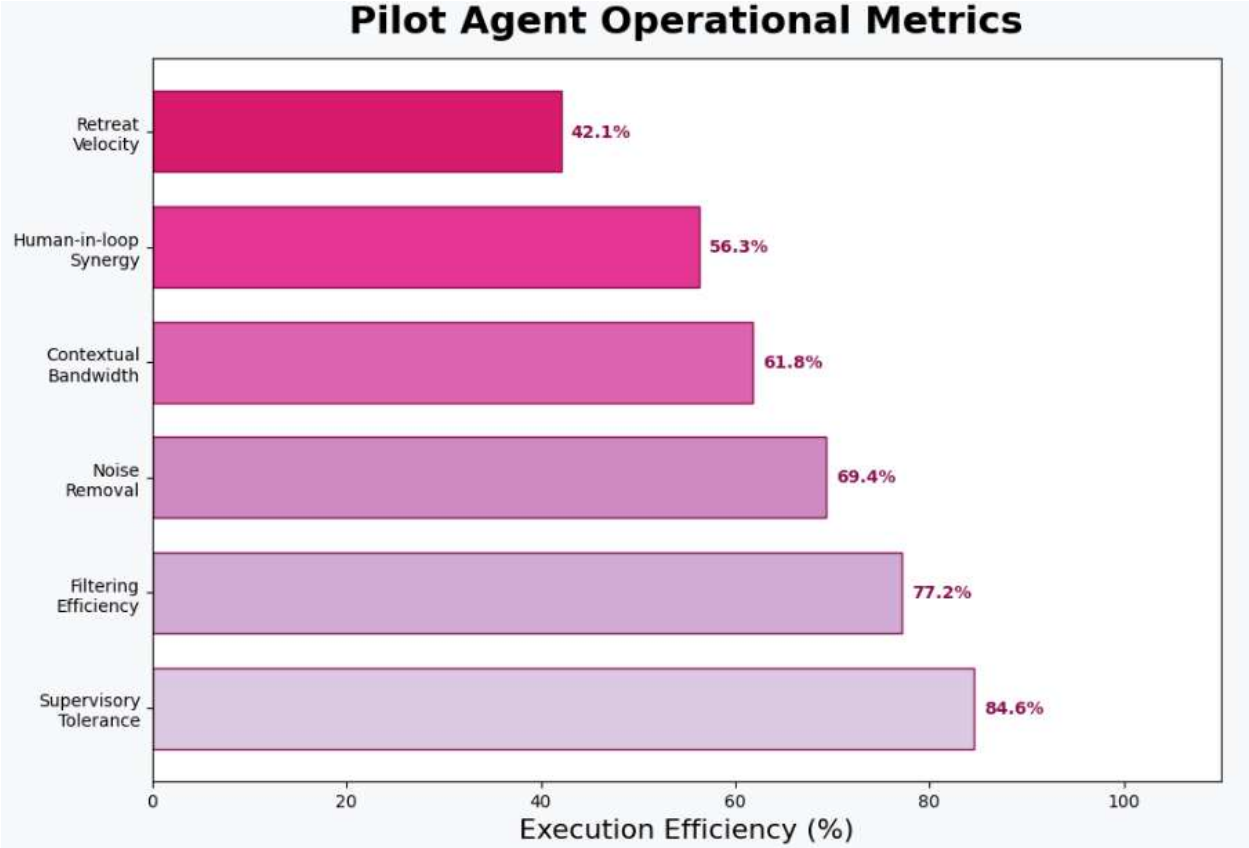


Fig 4: Pilot Agent Operational Metrics

12.1. Staged Deployment Strategies

The phased introduction of autonomous AI agents must address technical, operational, and human factors. Roadmaps typically outline three stages: pilot, scale, and rollback. Pilot deployment tests technical feasibility, operational readiness, and service levels; success indicators confirm readiness for scaled rollout. Digital banks should determine pilot scope, evaluate risks, define KPIs, and engage stakeholders to promote buy-in and support.



Pilot outcomes inform strategies for scaling the capability and transitioning pilot-mode support services into BAU. Factors such as application complexity, level of human-in-the-loop support, operational exposure, transaction volume, and consequent impact on operational resilience affect the time frame for scaling from pilot to full-speed operation. Prioritization strategies for staged scaling typically include risk, return, difficulty, and testing time. Service levels, in terms of availability, throughput, and latency, further influence pilot scaling decisions.

Although designed to operate automatically, autonomous AI agents can produce unexpected or undesirable outcomes. Systems and a change management process for controlling the scope, cadence, and duration of automated operation should therefore be implemented before first deploying an autonomous agent at scale. Such operations close a key-risk gap associated with autonomous operation, enabling a measured and reversible approach to scaling early-stage deployed models into production-grade capacity.

12.2. Change Management and Workforce Impacts

Change management frameworks support structured transitions, aligning stakeholders and addressing project stages. Ethnographic studies detail change preferences in banking operations, providing insights on workforce adaptation to emerging technologies. Autonomous agents shift operational models, complementing or replacing traditional roles. AI systems undertake repetitive, error-prone tasks, enhancing efficiency and freeing staff for strategic functions. These investments boost productivity and elevate competition. Workforce familiarity with AI nurtures utilization interest and positive perceptions, enabling change management and improving user experience.

Business processes driven by autonomous agents instead of human resources permit significant operational change through stakeholder input. Employees desire a share of savings and reinforced job security during difficult transitions. In the financial sector, autonomous agents offer a new platform with change specifications. Costly agents must be justified in simple, well-defined processing activities. The entire transaction process, including antimoney laundering compliance through identity verification, transaction monitoring, and customer support, falls within an extremely regulated environment. Productivity can be improved through the reduction of response times and error rates along the entire process. Savings can be enjoyed thanks to the reduction of workload in these tedious tasks. Employees can focus their expertise in really



complicated problems, where a critical thought must be put, leaving the details of a transaction supervision process to an agent.

13. Conclusion

Regulatory authorities and industry experts are calling for greater use of technology in the financial services sector through innovation that expands the use of data in ways that enhance customer experience, reduce cost to serve, and create deeper insights into the behaviour of customers. Generating new intelligence or predictive scoring based upon a large quantity of data, the majority of which have not been previously leveraged, would help the industry achieve these objectives.

The analysis is based upon applied research into contemporary banking systems architecture as part of the creation of an open-source platform for the finance domain. It contributes to the work-in-progress through the examination of end-to-end banking workflows—customer onboarding, including identity verification for both individuals and companies, and transactions processing with fraud detection and prevention—and their implications for the autonomy levels of embedded AI agents with a focus on the decision-making autonomy and the governance and control mechanisms.

References

1. Atwal, G., & Bryson, D. (2021). Antecedents of intention to adopt artificial intelligence services by consumers in personal financial investing. *Strategic Change*, 30(3), 293–298.
2. Seiler, V., & Fanenbruck, K. M. (2021). Acceptance of digital investment solutions: The case of robo advisory in Germany. *Research in International Business and Finance*, 58, 101490.
3. Kalyani, S., & Gupta, N. (2023). Is artificial intelligence and machine learning changing the ways of banking: A systematic literature review and meta analysis. *Discover Artificial Intelligence*, 3, 41.
4. Mangalampalli, B. M. Generative AI Applications In Healthcare Data Mart Design And Optimization.
5. Weber, P., Carl, K. V., & Hinz, O. (2023). Applications of explainable artificial intelligence in finance—A systematic review of finance, information systems, and computer science literature. *Management Review Quarterly*, 74, 867–907.



6. Schick, T., Dwivedi-Yu, J., Dessì, R., Raileanu, R., Lomeli, M., Hambro, E., Zettlemoyer, L., Cancedda, N., & Scialom, T. (2023). Toolformer: Language models can teach themselves to use tools. *Advances in Neural Information Processing Systems*, 36.
7. Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing reasoning and acting in language models. *International Conference on Learning Representations*.
8. Kolla, S. H. (2022). Strategic Information Integration Models for Cross-Functional Service Optimization in Large-Scale Enterprises. *International Journal of Emerging Trends in Engineering and Management Research*, 7(3), 11811.
9. Wu, Q., Bansal, G., Zhang, J., Wu, Y., Li, B., Zhu, E., Jiang, L., Zhang, X., Liu, S., Awadallah, A. H., White, R. W., Burger, D., & Wang, C. (2023). AutoGen: Enabling next-gen LLM applications via multi-agent conversation. *arXiv preprint arXiv:2308.08155*.
10. Park, J. S., O'Brien, J. C., Cai, C. J., Morris, M. R., Liang, P., & Bernstein, M. S. (2023). Generative agents: Interactive simulacra of human behavior. *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*, 1–22.
11. Shinn, N., Cassano, F., Gopinath, A., Narasimhan, K., & Yao, S. (2023). Reflexion: Language agents with verbal reinforcement learning. *Advances in Neural Information Processing Systems*, 36.
12. Amistapuram, K. (2023). Privacy-Preserving Machine Learning Models for Sensitive Customer Data in Insurance Systems. *Educational Administration: Theory and Practice*, 29(4), 5950-5958.
13. Li, Y., Wang, S., Ding, H., & Chen, H. (2023). Large language models in finance: A survey. *Proceedings of the Fourth ACM International Conference on AI in Finance*, 374–382.
14. Wu, S., Irsoy, O., Lu, S., Dabrovolski, V., Dredze, M., Gehrmann, S., Kambadur, P., Rosenberg, D., & Mann, G. (2023). BloombergGPT: A large language model for finance. *arXiv preprint arXiv:2303.17564*.
15. Malibari, N., Katib, I., & Mehmood, R. (2023). Systematic review on reinforcement learning in the field of FinTech. *arXiv preprint arXiv:2305.07466*.
16. Lakkaraju, K., Vuruma, S. K. R., Pallagani, V., Muppasani, B., & Srivastava, B. (2023). Can LLMs be good financial advisors? An initial study in personal decision making for optimized outcomes. *arXiv preprint arXiv:2307.07422*.
17. Gai, K., Qiu, M., & Sun, X. (2020). A survey on FinTech. *Journal of Network and Computer Applications*, 187, 103173.
18. Yandamuri, U. S. (2023). An Intelligent Analytics Framework Combining Big Data and Machine Learning for Business Forecasting. *Zenodo*.



19. Ozili, P. K. (2020). Financial inclusion and FinTech during COVID-19 crisis: Policy solutions. SSRN Electronic Journal.
20. Frost, J. (2020). The economic forces driving FinTech adoption across countries. BIS Working Papers, 838.
21. Boot, A., Hoffmann, P., Laeven, L., & Ratnovski, L. (2021). FinTech: What's old, what's new? Journal of Financial Intermediation, 48, 100945.
22. Goldstein, I., Jiang, W., & Karolyi, G. A. (2021). To FinTech and beyond. The Review of Financial Studies, 34(5), 1647–1669.
23. Inala, R. Designing Scalable Technology Architectures for Customer Data in Group Insurance and Investment Platforms.
24. Thakor, A. V. (2020). FinTech and banking: What do we know? Journal of Financial Intermediation, 41, 100833.
25. Fuster, A., Goldsmith-Pinkham, P., Ramadorai, T., & Walther, A. (2022). Predictably unequal? The effects of machine learning on credit markets. The Journal of Finance, 77(1), 5–47.
26. Bussmann, N., Giudici, P., Marinelli, D., & Papenbrock, J. (2021). Explainable machine learning in credit risk management. Computational Economics, 57, 203–216.
27. Arner, D. W., Barberis, J., & Buckley, R. P. (2020). The evolution of FinTech: A new post-crisis paradigm? Georgetown Journal of International Law, 47(4), 1271–1319.
28. Weber, M., Domeniconi, G., Chen, J., Weidele, D. K. I., Bellei, C., Robinson, T., & Nalisnick, E. (2020). Anti-money laundering in Bitcoin: Experimenting with graph convolutional networks for financial forensics. arXiv preprint arXiv:1908.02591.
29. KollIntegration. South Eastern European Journal of Public Health, 248–260.
30. Hildebrandt, T., Sawaya, G., & Kearney, C. (2022). Artificial intelligence and machine learning in banking: Opportunities and risks. Journal of Financial Transformation, 55, 77–89.
31. Tam, K., & Jones, K. (2021). Machine learning and artificial intelligence in financial services: Applications, risks, and governance. Journal of Financial Regulation and Compliance, 29(4), 421–438.
32. Choudhury, A., & Bhowmik, R. (2022). Artificial intelligence in banking: Applications, challenges, and future directions. International Journal of Financial Engineering, 9(3), 2250021.
33. Jagtiani, J., & Lemieux, C. (2021). The roles of alternative data and machine learning in fintech lending: Evidence from the LendingClub consumer platform. Financial Management, 50(4), 1007–1036.



34. Bracke, P., Datta, A., Jung, C., & Sen, S. (2020). Machine learning explainability in finance: An application to peer-to-peer lending. Staff Working Paper, Bank of England, 816.
35. a, S. K., & Reddy, V. A. R. (2023). Deep Learning Architectures For Multimodal Medical Data
36. Kalyani, S., & Gupta, N. (2022). Artificial intelligence and machine learning applications in banking: A systematic review of emerging research. *International Journal of Management and Economics*, 58(4), 345–362.
37. Oliveira, T., Thomas, M., Baptista, G., & Campos, F. (2021). Mobile payment: Understanding the determinants of customer adoption and intention to recommend the technology. *Computers in Human Behavior*, 61, 404–414.