



sumsub

SFA

SINGAPORE
FINTECH
ASSOCIATION

Sumsub APAC State of Digital Trust: AI Governance Benchmark

A guide to assessing Autonomy,
Responsibility, and Traceability
in AI across organizations
in Asia-Pacific

Table of contents

Foreword	03
The AI adoption trend	05
Methodology	08
The Accountability Asymmetry	12
Other observations	19
Breakdown by sector	28
Breakdown by market	34
Regulatory anchors at a glance	42
What comes next	44
AI governance readiness checklist	47
How Sumsub can help	50

Foreword

For the past few years, most of us asked one question about AI: what can it do for us? A weightier question has taken its place: what is it doing on our behalf, and can we prove it?

AI has moved past the role of an assistant. Systems that once drafted a memo or flagged an anomaly now open tickets, move money, approve transactions, and complete multi-step tasks across the business with limited human input. These are AI agents: software that pursues a goal, makes choices along the way, and takes actions inside live systems on its own. AI agents are starting to behave like participants in the business, and increasingly they transact with the world beyond it.

This change reshapes what AI governance has to deliver. When a person makes a decision, accountability is clear. **When an AI agent makes thousands of decisions a day, the organization has to answer three questions on demand: what did the AI do, why did it do it, and who is answerable for the outcome?**

This guide measures how ready organizations across APAC are to answer those three questions.

The AI adoption trend

Across APAC, AI has moved past isolated pilots in most organizations. Nearly three-quarters use it across multiple functions. It runs most widely in analytical and operational work such as data analysis, operations, and customer service, though use cases split by sector: more than half of Financial Services businesses use it for compliance and case review, and almost three in four use it for fraud detection.

General adoption is no longer the interesting part; AI increasingly acts on its own, running workflows, making calls, and carrying tasks through several steps without a person in the loop for each one.



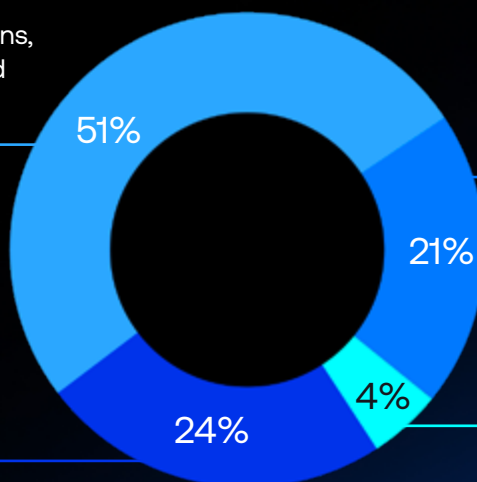
Chart 1: The role AI currently plays in your organisation

Execute with approval

Handles some decisions or actions, with human involvement required for key steps or approvals

Automate

Automates routine tasks while keeping human approval of outcomes



Autonomous agent

Mainly supports staff with information or recommendations

Support and recommend

Handles some decisions or actions independently, with human involvement only for exceptions or oversight

From the survey, 72% of organizations report that AI either handles decisions and actions with approval (51%) or acts independently with exception-based oversight (21%), usually for internal teams or customers and less often with external business stakeholders. Human oversight thins out as you move down the curve.

The pressing question now is how to govern AI that acts for us, and that is what this benchmark assesses across three pillars: **Autonomy, Responsibility, and Traceability.**

Methodology

This study was conducted independently by Blackbox Research, commissioned by Sumsu in partnership with the Singapore FinTech Association. A roundtable with selected participants from Crypto and Financial services framed the research and surfaced trends on the ground.

A quantitative survey then reached 720 senior professionals in technology, product, risk, compliance, and operations roles across nine APAC markets and four sectors. The study took place over three months, between April and June 2026.

The benchmark scores each organization from 0 to 100 across **three pillars** using a dedicated set of survey questions. A higher score means more advanced governance.

Autonomy

How far AI already acts on its own inside the organization

Whether AI executes multi-step tasks with limited human input

The most advanced role AI plays

The human involvement required when it acts

The breadth of autonomous agent roles

The expected trajectory of AI autonomy

Responsibility

How clearly the organization owns and answers for what its AI does

Who is accountable for AI-driven outcomes

Who is held responsible when AI gets it wrong

The governance controls in place

The level of approval required before agents act

Traceability

Whether the organization can reconstruct, explain, and audit what its AI did

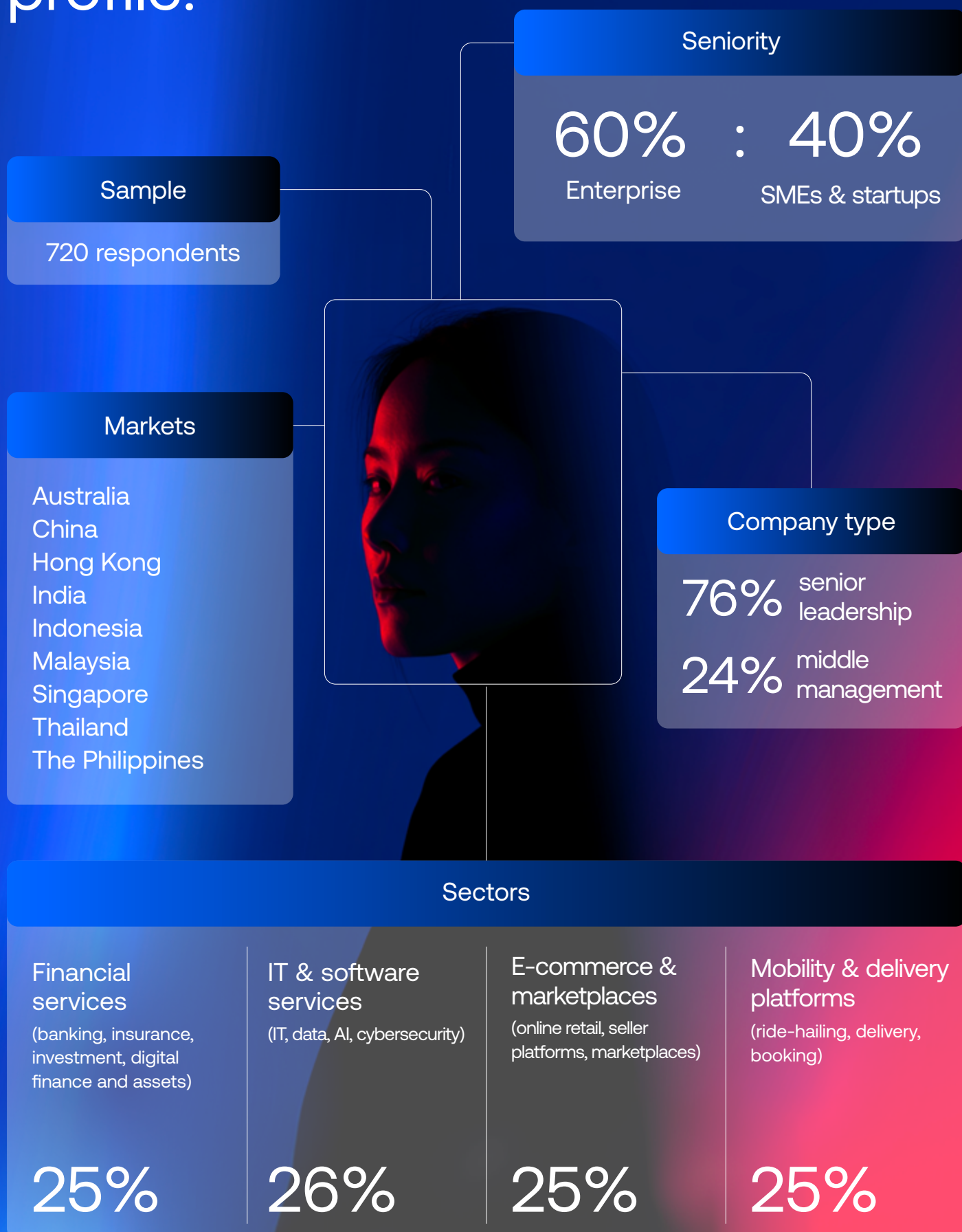
Confidence in tracing and explaining AI decisions

The capabilities in place to support traceability

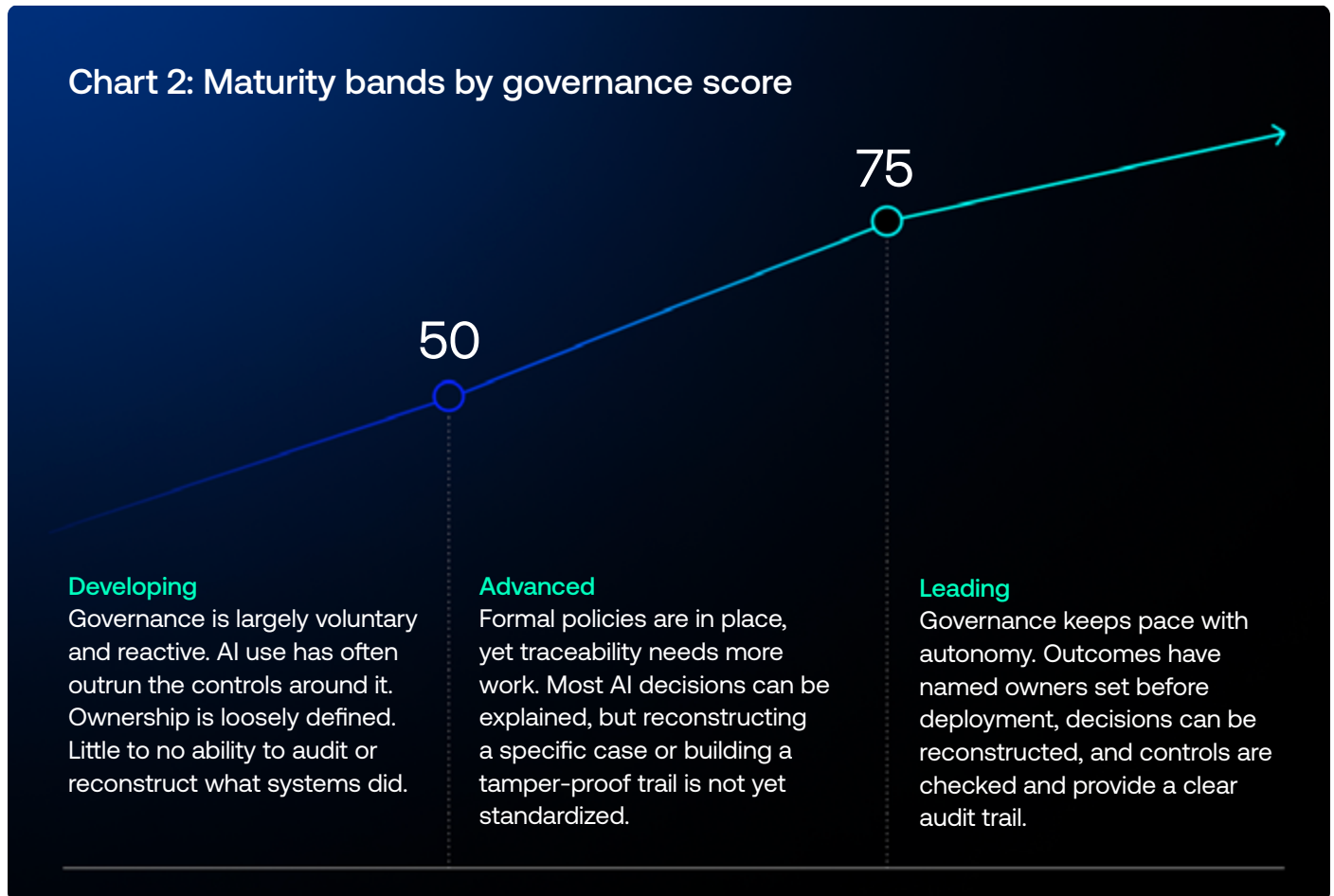
The gaps and barriers that remain

Awareness of the risks if they are not closed

Respondent profile:



The study also sorts organizations into three maturity bands by governance score:



The Accountability Asymmetry

Traceability lags behind Autonomy and Responsibility in every sector and every market. APAC organizations score 67.1 out of 100 on overall AI governance, but the headline number hides the pattern that matters. Autonomy and Responsibility both sit close to 70, while Traceability drops to 61.0.

We call that gap the Accountability Asymmetry: the distance between how fully organizations own their AI's decisions and how well they can prove what those decisions were. Autonomy multiplies it: a responsibility you've accepted but can't reconstruct, or audit becomes a liability you can't quantify the moment a regulator, shareholder, or customer asks you to show what happened.

Chart 3: AI governance score in APAC



The gap is really about proof

Organizations across APAC are stepping up to own AI-driven outcomes. What they lack is the ability to produce the evidence. 95% of respondents say they are confident they can explain an AI decision, yet only 50% can reconstruct the decision pathway, and just 38% hold a tamper-proof audit trail. The real challenge is whether those decisions hold up when a regulator, a shareholder, or a customer probes further.

Autonomy is outrunning control

APAC has entered a high-autonomy phase, with systems executing decisions at a speed and scale beyond direct human oversight. Traceability is not keeping pace. That creates a structural risk: enterprises are scaling decision-making without the infrastructure to monitor, audit, or step in, which widens exposure across compliance, operations, and customer trust.

A useful way to picture it

Autonomy is the gas pedal, responsibility is the driver insisting they are in control, and traceability is the steering wheel. Right now, APAC enterprises are pressing the gas without a steering wheel in their hands. Many are expanding their use of AI, many are willing to own its decisions, and many still cannot trace what that AI did or correct it after the fact.

Autonomy works as a multiplier here. When AI makes decisions at a volume and speed no team can track by hand, the gap between what it does and what can be reconstructed only widens. The more pronounced that gap grows, the higher the eventual cost, because by then the AI has executed thousands of actions no one can account for.



Mick Amelishko
Senior Engineering
Manager at Sumsub

“Every action an AI takes without a traceable record is a cost deferred, not avoided. Explainability is what lets you see the bill before it comes due.”

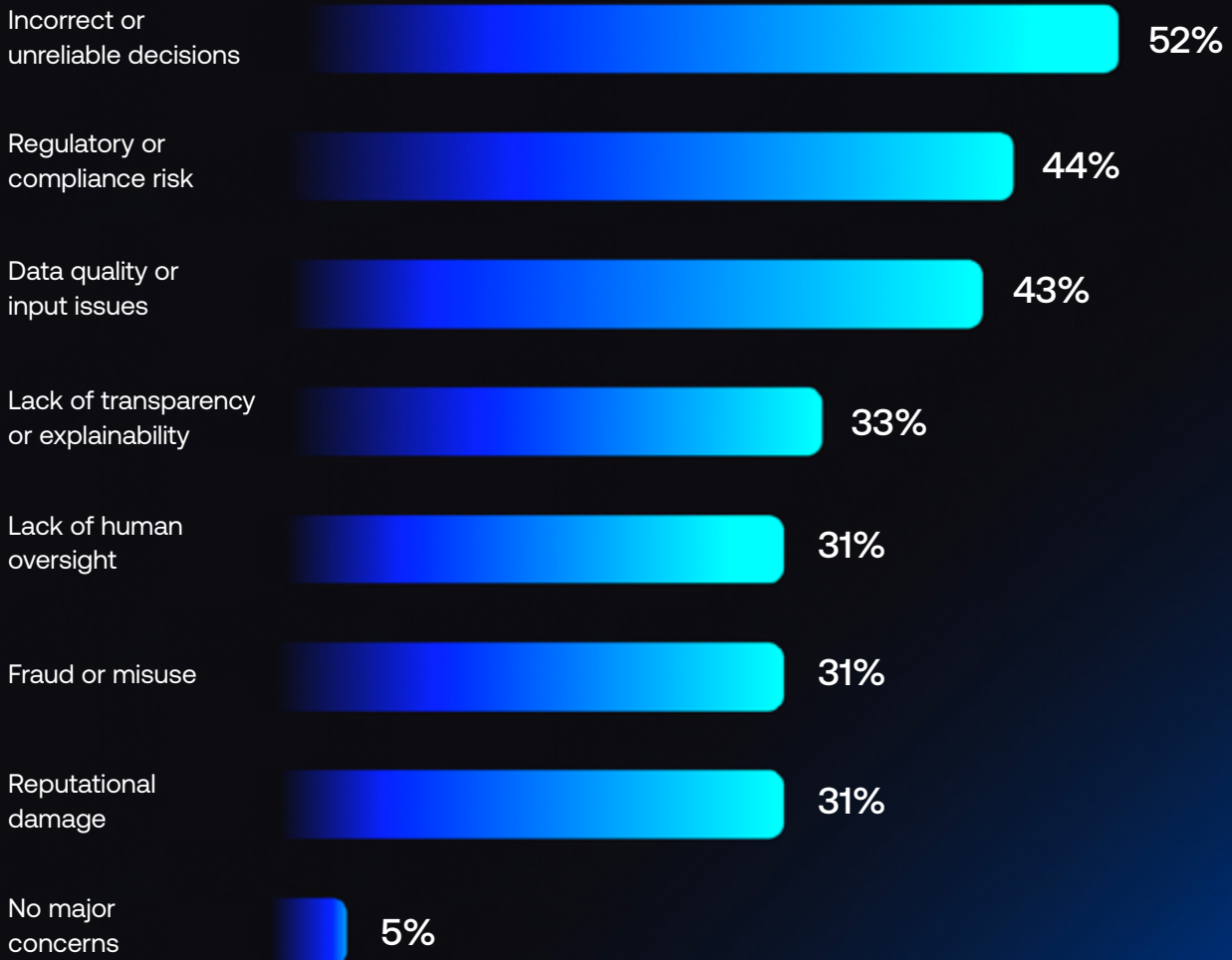
A spectrum of autonomy

“AI Agent” is a catch-all phrase. The label now covers a wide spectrum. Where a system sits on that spectrum decides how much can happen without a person in the loop. The table below outlines the scope of responsibility: the further right a system is on it, the thinner the oversight and the greater the Accountability Asymmetry becomes, because actions accumulate faster than anyone can reconstruct them.

Aspect	Assistant	Co-pilot	Agent
Autonomy: How much human interaction is needed?	Low autonomy Reactive only; acts when instructed in Q&A fashion	Medium autonomy Can automate parts of a task and chain a few calls, but under human supervision	High autonomy Acts independently after initial goal, hides intermediate steps from a user
Scope of tasks: How complex is the task the AI can handle in a single interaction?	Narrow Concrete intent from a prompt	Moderate Multi-steps task under supervision	Wide Manages entire workflows or objectives; coordinates many steps/tools (broad scope)
Interaction type: How does the AI interact with the user?	On-demand One prompt, one action (answer)	Continuous Collaboration; integrated into the user's workflow with ongoing suggestions. Often integrated into the work environment	Minimal Ongoing interaction; works in the background after kickoff, reports outcomes, or asks only if needed
Proactiveness: Does the AI initiate or adapt its actions?	None Does not initiate or suggest; waits passively for instructions	Semi-proactive Anticipates needs or suggests next steps in context, as the user progresses through a task	Fully proactive Pursues goals independently, adapts plan and actions based on new data or outcomes, without prompting

Chart 4: What worries organizations most when AI acts on their behalf

Respondents could choose multiple answers for this question.



The concerns that dominate are practical ones: whether the decision can be trusted and reconstructed. Incorrect or unreliable decisions (52%), data quality or input issues (43%), and a lack of transparency or explainability (33%) are reliability-and-evidence problems, and regulatory or compliance risk (44%) sits right alongside them. Fraud or misuse, the headline-grabbing version of AI risk, ranks lower at 31%.

With only **5%** reporting no major concern, the signal is close to unanimous: **95% of organizations see real exposure when AI acts for them**, and the exposure they name is operational and evidentiary, the everyday question of whether the output holds up under scrutiny.

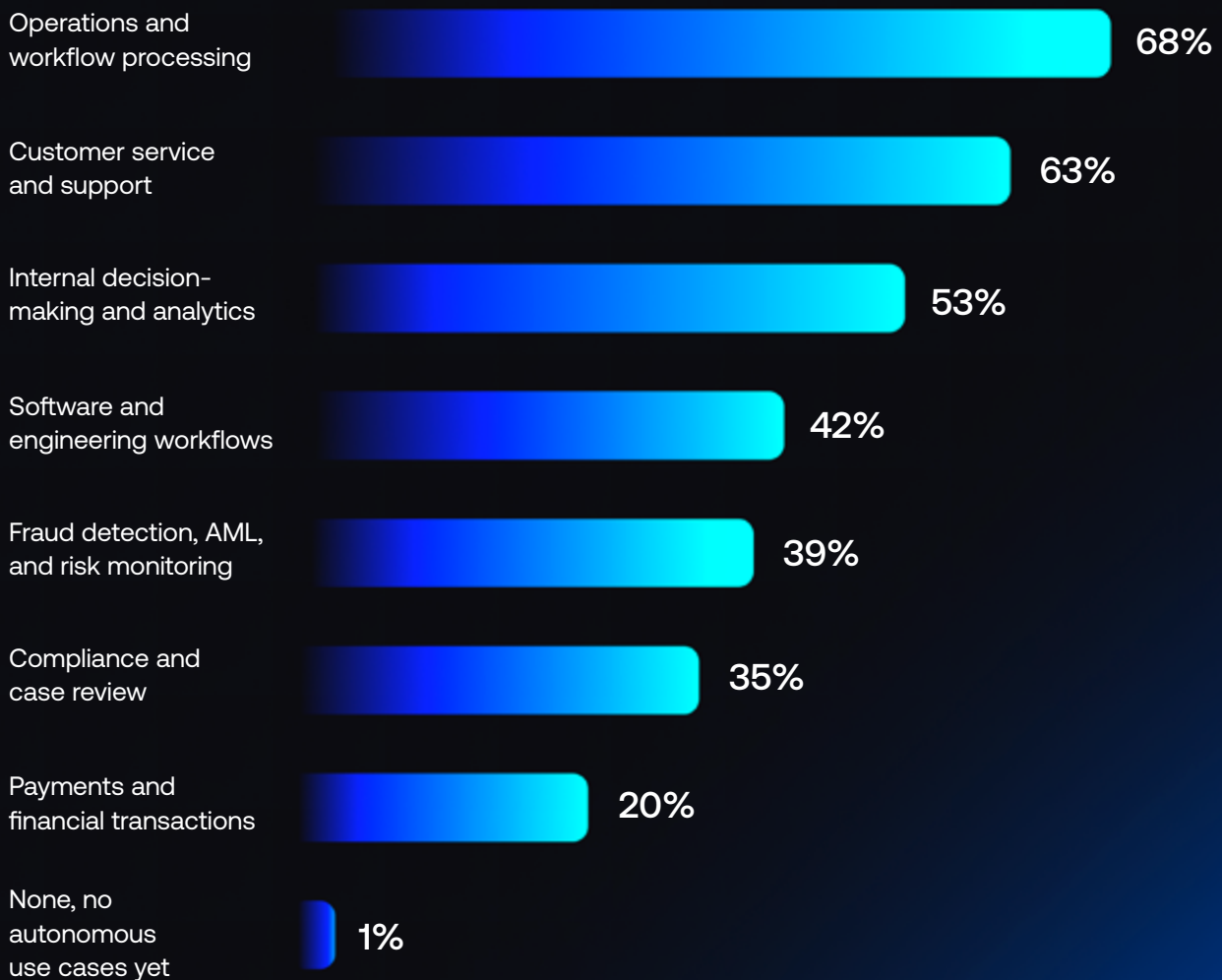


Other observations

AI turns agentic first where the stakes are lowest. Autonomous agents show up first in operations, customer service, and internal analytics. Payments, compliance, and anti-money laundering (AML) all trail behind because these areas are where a mistake costs the most. There is more caution with the adoption of new processes.

Chart 5: Where AI is taking on autonomous agent roles

Respondents could choose multiple answers for this question





A misrouted ticket costs only a few minutes. Fraud detection, AML, and risk monitoring look like the natural home for agents (pattern analysis, rule-following, triage, hand-off to the next step), yet only 39% report autonomous agent roles there and just 35% in compliance and case review. On paper, these are close to an ideal use case, yet the caution is understandable: agentic AI is still a novel technology that requires careful adoption and supervision.



Holly Fang
President of the
Singapore FinTech
Association

“The path to agentic AI will likely be incremental rather than binary. Fintechs are adopting AI agents first in lower-risk workflows, while retaining human oversight for decisions involving customer onboarding, the movement of client funds, and other regulated activities. Building confidence in governance will be just as important as advances in the technology itself.”

Three things hold it back:

- 1 The cost of a wrong decision is asymmetric and severe
- 2 A missed sanctioned entity or an unfiled suspicious activity report carries fines, license risk, and personal liability for named officers
- 3 A false positive drags a legitimate customer through friction and possible fair-treatment complaints.

With only 50% able to reconstruct a decision pathway and 38% holding a tamper-proof audit trail, handing the highest-stakes calls to an agent means scaling actions no one can explain on demand. This tension is especially pronounced in regulated domains, such as anti-money laundering, where accountability structures are already well-defined. Accountability in AML is anchored in a person: a money-laundering reporting officer signs off, and a regulator expects a human to answer.

The lag is the Traceability gap doing its damage precisely where the stakes are highest. An agent can only be trusted in fraud prevention and AML work if the firm can reconstruct and defend every call it makes when a regulator asks, which makes the audit trail the gating requirement. Financial services is better placed than any other sector, with 68% keeping audit trails of AI decisions, the highest anywhere. Even there, the harder capability is rare: only 41% hold secure, tamper-proof records, level with IT and ahead of the rest, but well short of what regulator-grade defensibility demands once agents are making these calls at scale.



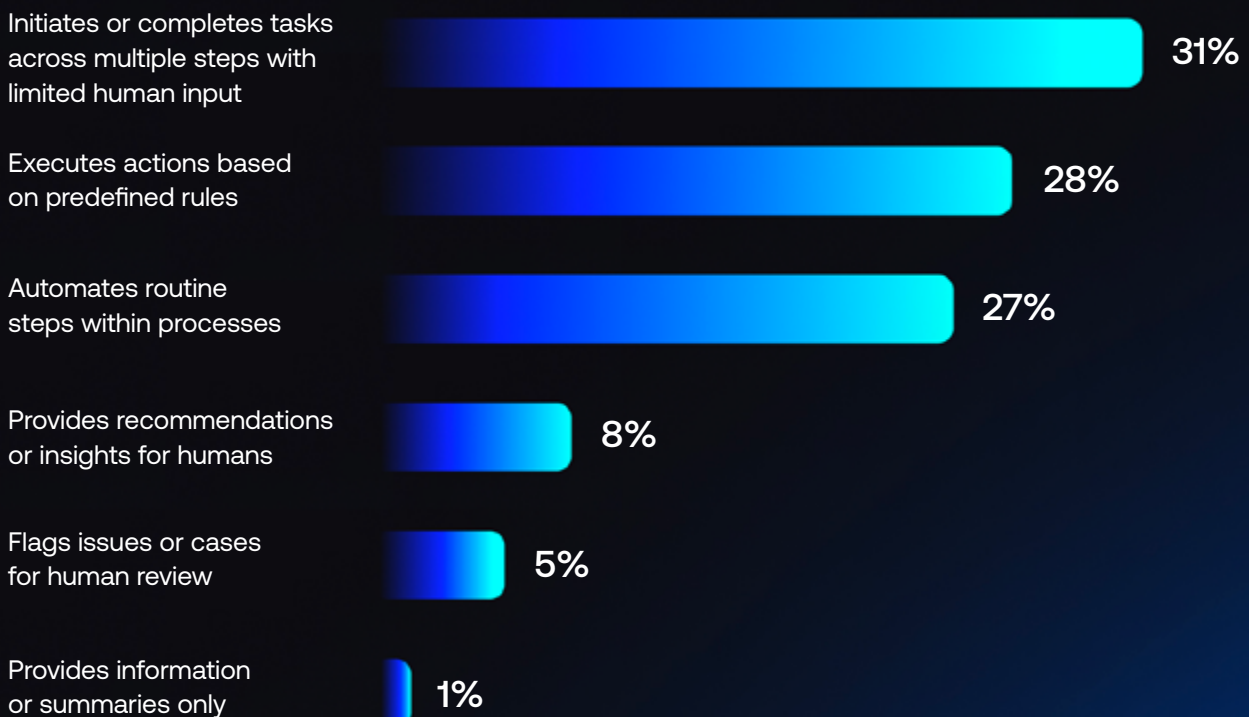
Nimit Gulati
Vice President,
Acceptance Solutions,
Value-Added Services at
Visa, Asia Pacific

“The real challenge is not whether AI can make decisions, but whether those decisions can be trusted. This research suggests that organizations are moving more cautiously when it comes to payments and financial transactions than in other business functions, reflecting the higher bar for accountability when money is involved. The findings highlight the importance of transparency, security and accountability in AI-driven commerce.”

True agentic AI is still in its infancy

Only 31% say AI initiates or completes multi-step tasks with limited human input, the agentic end of the curve. For most organizations, rule-based execution (28%) or routine automation (27%) is still the ceiling.

Chart 6: Most advanced role played by AI in your organization



Most enterprises still hit an AI deployment bottleneck and cannot capture the full cost savings AI promises. The divide underneath is foundational: true agentic AI needs unstructured data access, cross-platform APIs, and fluid software integration. Organizations stuck in routine automation are more often held back by legacy technical debt, siloed databases, and rigid architectures. The 31% at the agentic end tend to have cleaner data pipelines and stronger digital infrastructure, so the distance between them and everyone else is an infrastructure gap that takes years and real capital to close.

AI use is outgrowing its controls

More than nine in ten respondents admit they have stretched an AI system past its original purpose in the past year, and roughly six in ten have already had an autonomous AI action go wrong.

Chart 7: Expanded an AI system beyond its intended use in the past 12 months

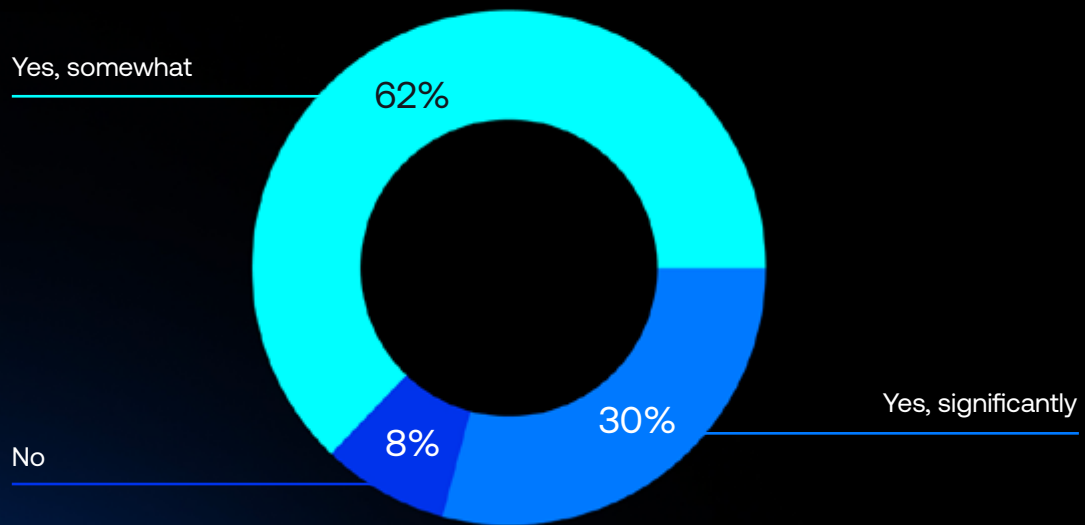
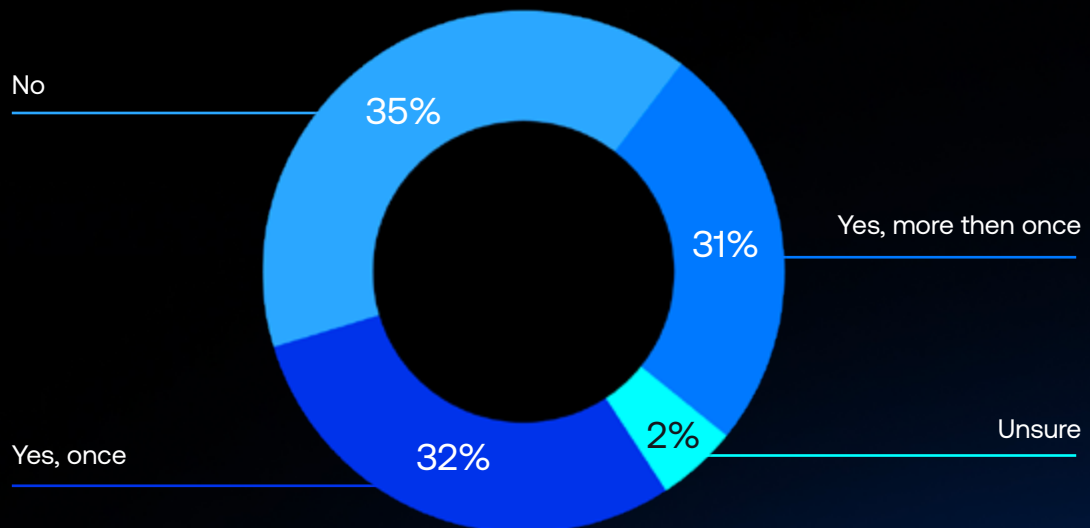


Chart 8: Experienced an autonomous AI action with an unintended or problematic outcome





92% have widened a system's scope in the past year, with 30% doing so significantly, as the AI running today is rarely the AI that was originally reviewed and signed off. 63% have already watched an autonomous action produce an unintended or problematic outcome, and 31% have seen it happen more than once.

Scope is expanding faster than the controls around it, which is the Accountability Asymmetry playing out in real time: more autonomy, more drift, and a growing stack of actions that were never governed for the job they ended up doing.

Breakdown by sector

Financial services (40%) and IT & software services (39%) are nearly twice as likely as E-commerce and marketplaces (21%) or Mobility and delivery platforms (23%) to run genuinely agentic AI, where AI initiates or completes multi-step tasks with limited human input.

Chart 9: Breakdown by sector





Edmund Phuang
Digital Data & AI
Adoption Lead,
CIMB Singapore

“One of the key values of large language models is their ability to institutionalize knowledge. For example, banks may process hundreds of credit memos, but any individual typically gains exposure to only a fraction of them. By capturing that collective expertise, new joiners can start from the bank’s accumulated best practice rather than from scratch. AI can create the first draft or recommendation, but the human remains accountable for reviewing, validating and making the final decision.”



Financial services

The most mature sector, and the one where audit trails carry the most weight. Financial firms are moving fast, with 72% running multi-step AI systems in production, and they post the highest overall governance score at 69.6, leading on both Responsibility and Traceability.

Strict regulatory standards leave little choice: the biggest perceived risks when AI acts for the firm are regulatory or compliance risk (54%) and fraud or misuse (34%). If an automated system triggers an unintended payment, the absence of an immutable audit trail turns straight into liability. 68% keep audit trails of AI decisions, ahead of IT (59%), e-commerce (50%), and mobility (55%).



IT & software services

The most aggressive sector. 94% carry out multi-step processes autonomously, and 42% have pushed AI systems significantly past their original scope in the past year, the highest of any sector. It also holds the largest share of Leading organizations at 31%. That high-speed environment breeds agentic drift, where software workflows run semi-independently, leaving firms exposed to data quality issues (50%) and a lack of clear regulatory guidance (46%).



E-commerce and marketplaces

Moderate maturity, with a mean governance score of 65.4. This sector leans on practical, operational guardrails more than strategic governance structures. AI runs heavily in customer service (71%) and sales and marketing (69%), used as a utility to drive sales and handle complaints rather than as a core engineering asset, which fits a business built on thin margins and front-end experience. The imbalance shows up in the controls: 69% run a deployment approval process, but only 29% subject their systems to independent audit or validation.



Mobility and delivery platforms

Lowest across all four measures. Logistics, courier, and ride-hailing networks tend to treat algorithmic dispatching and routing as physical math rather than corporate decision-making, so governance frameworks get deprioritized. That is risky, because this AI physically triggers real-world delivery, driver payouts, and routes. A systemic glitch in an unmonitored, poorly traceable system turns into immediate field disruption, labor friction, and supply-chain delay. 41% here fear the loss of human oversight, and the sector holds the largest share of organizations still in the Developing band.

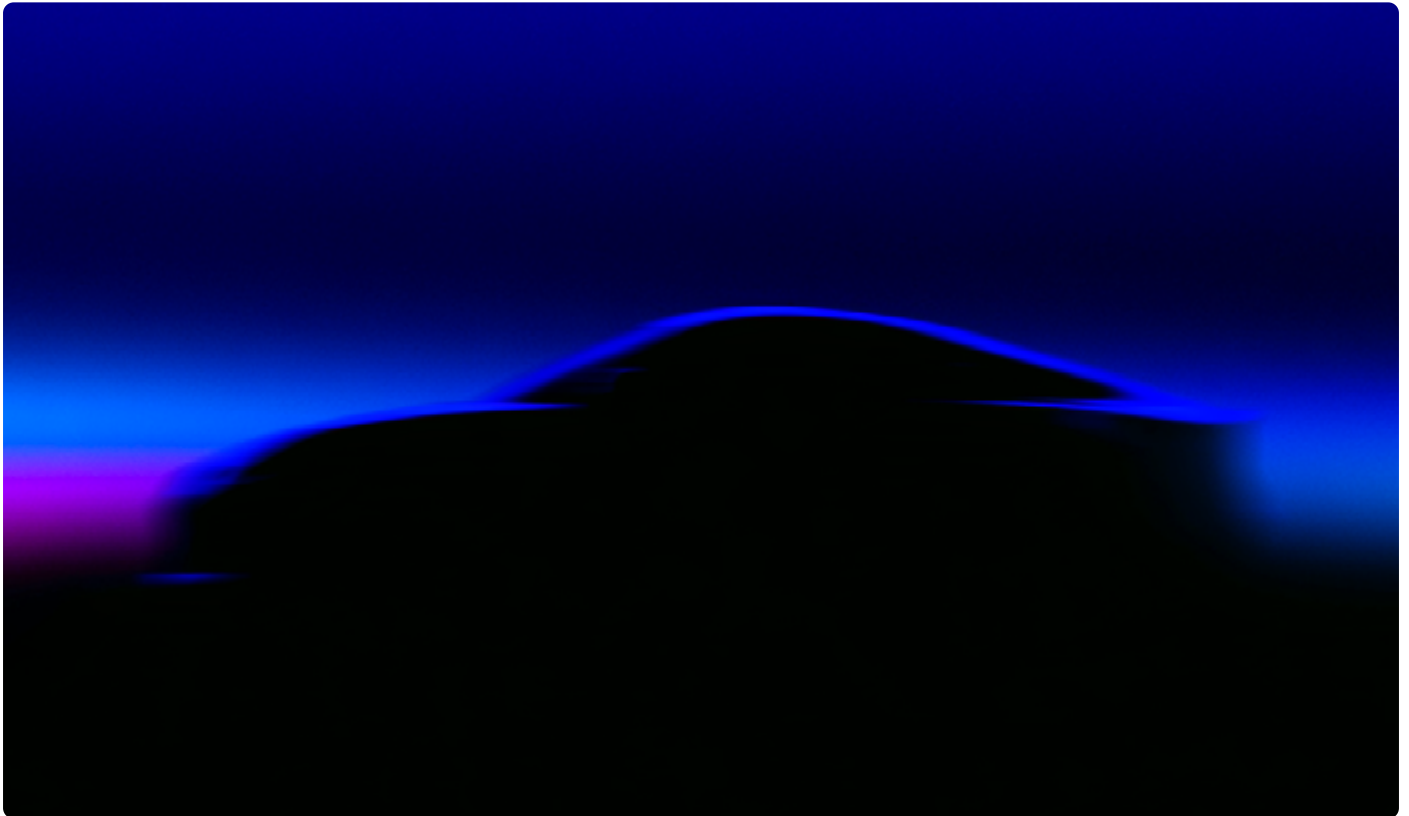
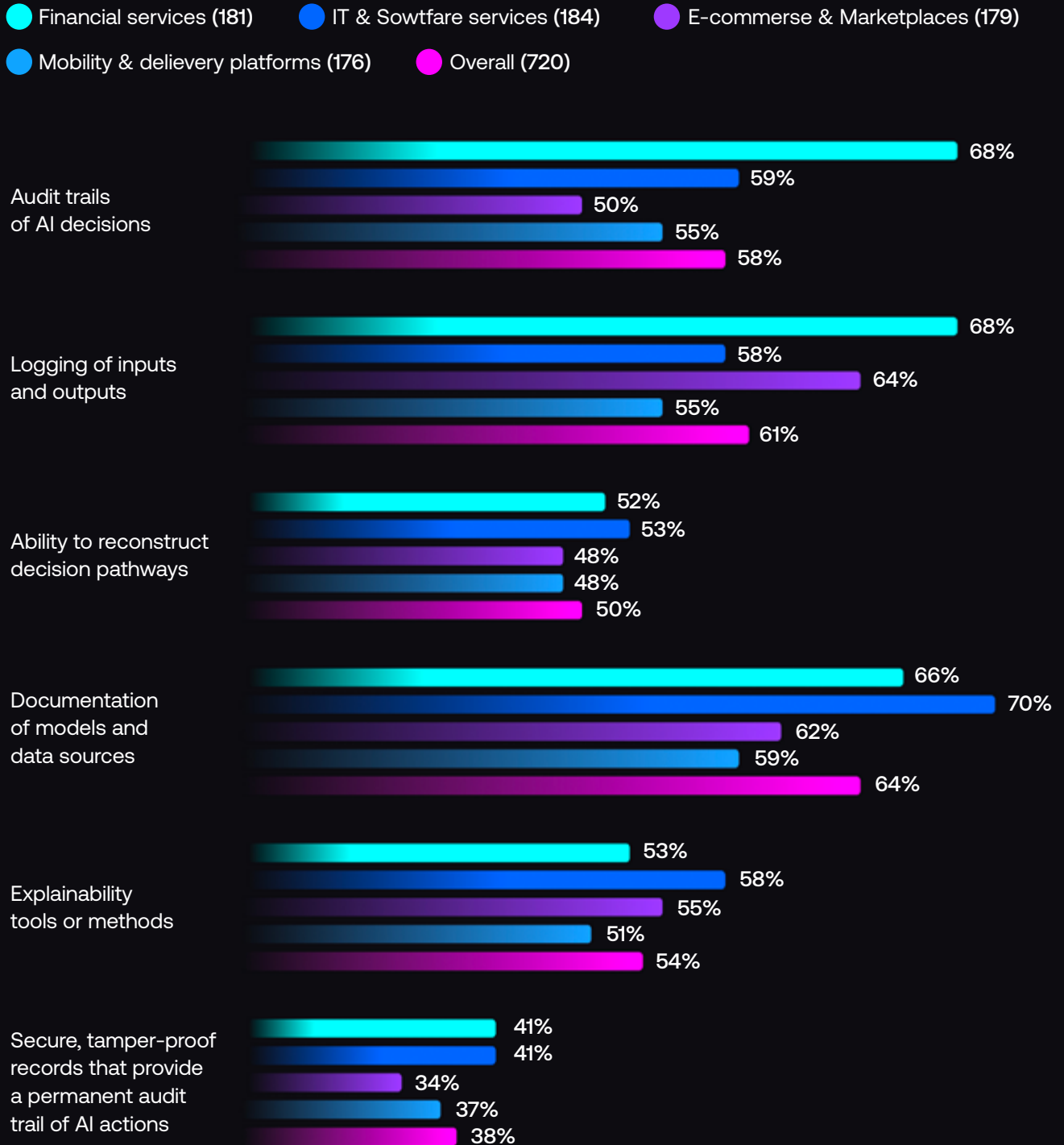


Chart 10: Capabilities that organizations currently have in place to support traceability of AI decisions

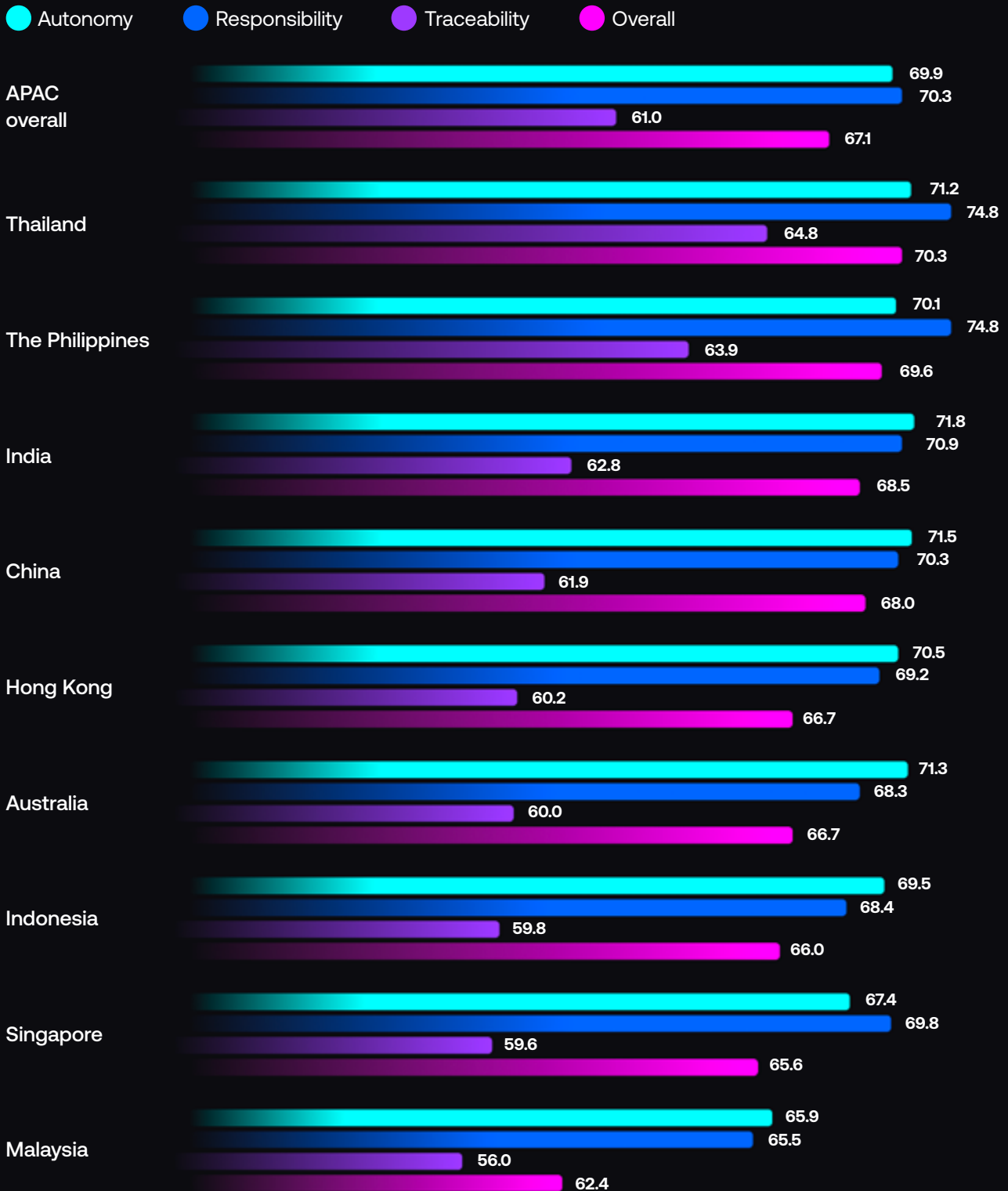
Respondents could choose multiple answers for this question.



Breakdown by market

Every market shows the same shape: Traceability sits below Autonomy and Responsibility. The spread between the top and bottom of the table is 7.9 points on the composite.

Chart 11: Breakdown by market



The leaders



Thailand tops the table at 70.3, with the region's best Traceability score and a Responsibility score of 74.8 that ties the Philippines for the highest of any market. Nearly a third of Thai organizations are rated Leading. It is building modern governance on greenfield digital foundations, backed by the Royal Decree on Digital Platform Services and national AI Ethics Guidelines.



The Philippines follows at 69.6, with 27.5% of organizations Leading and 73% keeping audit trails of AI decisions, the highest of any market. Its compliance-first posture is driven by its role as a global business process outsourcing (BPO) hub, where the need to win enterprise contracts and shield against data liability drives rigorous policy checklists aligned to standards like the EU AI Act.

India, and the confidence paradox



India ranks third at 68.5 and ties for the highest share of Leading organizations, on the back of elite, paper-based compliance built to meet international standards like the EU AI Act. Its national framework, the India AI Governance Guidelines from the Ministry of Electronics and IT, favors rapid deployment and self-regulation, and that appetite for speed shows up in production.

That is the paradox at the top of the table: Thailand and India both post strong governance scores, yet both also report the region's highest rate of active AI mistakes, at 68.8% of organizations. The confidence in what AI can do is running ahead of how reliably it performs once deployed.

China: a national ID for every AI agent



China reads differently from the rest of the table. It sits fourth at 68.0, but almost every surveyed company (roughly 97.5%) falls inside the Advanced or Leading bands, with only two rated Developing. That compression reflects a hard-regulation climate; the Cyberspace Administration of China runs algorithm registries, data localization rules, and pre-launch security reviews, backed by enforcement sweeps that penalize misconduct with revenue-based fines or service suspensions. Compliance is the price of market entry, so a disciplined baseline is the norm.

China is also the clearest signal of where agent governance is heading. In June 2026, its standards body issued a national standard for AI agent interconnection built around a unified identity framework, described in state media as a digital ID card for every AI agent, so agents can be recognized and trusted across systems. The logic maps directly onto this report's central finding: as agents start to act and transact, the way to manage trust is to bind each agent to a verifiable identity. That is the Know Your Agent (KYA) concept, and a national regulator has now written policy around it.

The middle and the tail

Both Hong Kong (66.7) and Australia (66.7) cluster heavily in the Advanced band and regulate AI through existing privacy and consumer law rather than a standalone AI act.



In Hong Kong, the Department of Justice has convened an Inter-Departmental Working Group to review where existing laws need to adapt to the fast, pervasive use of AI, tasking each bureau with scrutinizing the areas of law under its remit.



Australia is reacting to real liability ahead of the December 10, 2026 automated decision-making transparency amendments under the Privacy Act, with Australian Consumer Law penalties now doubled to AU\$100 million per breach.



Indonesia (66.0) sits in the middle of a shift from voluntary ethics guidance to binding rules. Its financial regulator, OJK, moved early with banking AI governance guidelines that mandate lifecycle risk management and model validation, and the government is working to elevate its national AI strategy into a binding Presidential Regulation.



Singapore (65.6) ranks low for a mature market, and that is the point. Its conservative self-grading is a mark of discipline and a mature governance culture. With initiatives such as the Model AI Governance Framework for Agentic AI and AI Verify developed by the Infocomm Media Development Authority, as well as the Safeguards for Agentic Finance at Runtime developed by the Monetary Authority of Singapore, firms in Singapore measure themselves against technical benchmarks rather than paper checklists and grade honestly.



Malaysia (62.4) sits last, carrying the region's highest concentration of Developing organizations. Its historically voluntary approach is about to end: a statutory AI Governance Bill with a mandatory, risk-based model and real penalties is finalized and heading to cabinet, which leaves many firms scrambling to build formal frameworks from scratch.

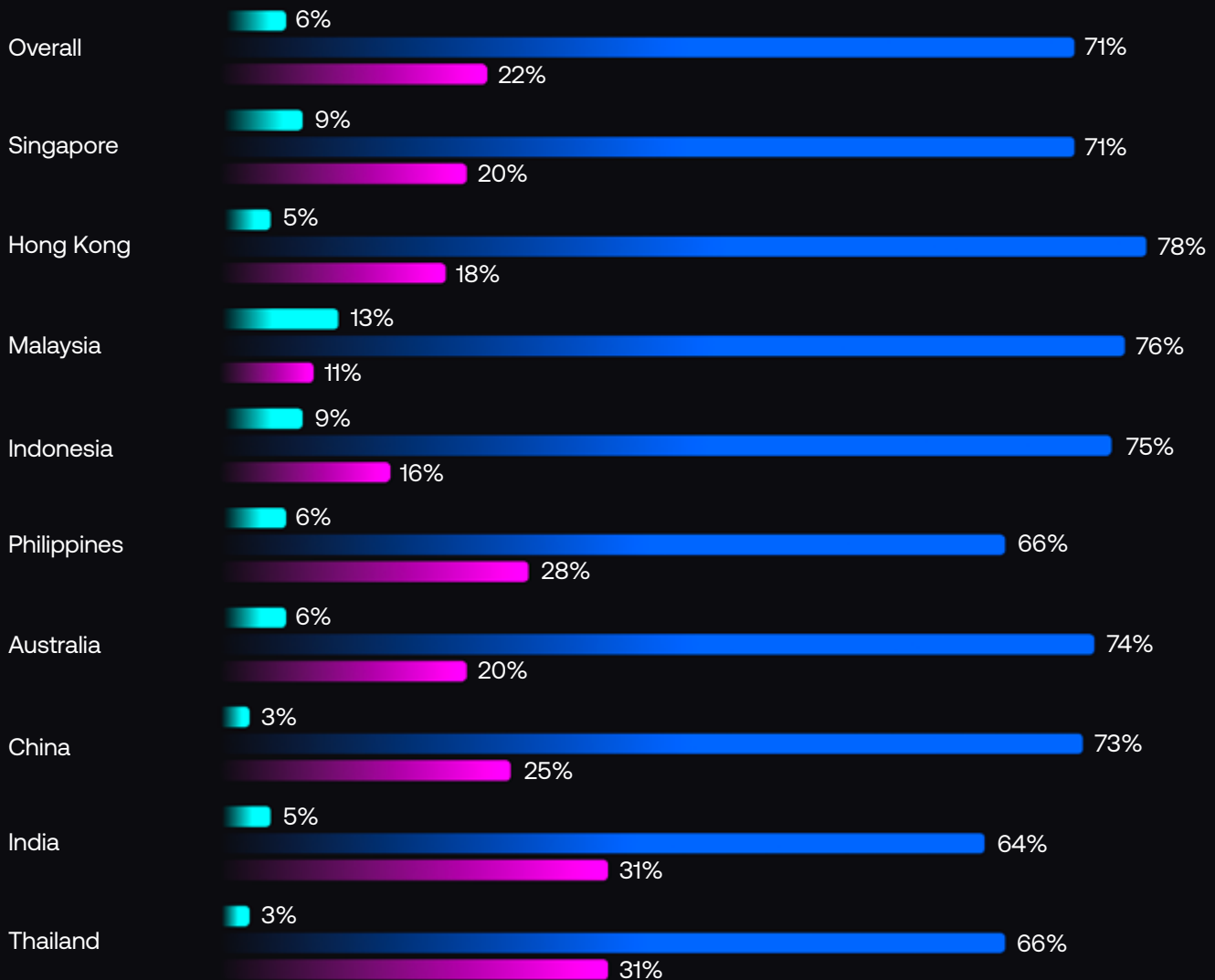


What ties this band together is regulation in transition. Each of these markets is either stretching existing privacy and consumer law to cover AI or waiting on new mandates still to take full effect. Their scores track frameworks that are still settling into practice. Singapore is the exception, where a low score signals rigor rather than a gap in capability.

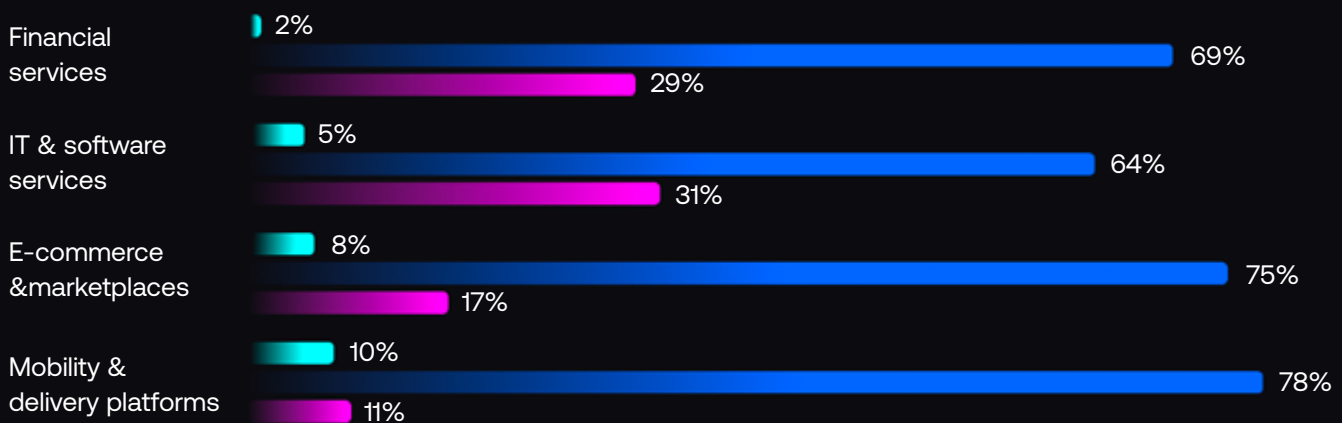
Chart 12: AI Governance maturity score

Developing Advanced Leading

Country:



Sector:





Penny Chai
Vice President,
APAC, SumsuB

“Prudence, rather than a lack of strategic intent, defines how enterprises are scaling AI agents. When financial liabilities are on the line, immature traceability systems create an unacceptable operational risk. Establishing robust tracking architectures and guardrails is the vital prerequisite to safely deploying high-stakes AI at scale.”

Regulatory anchors at a glance

The scores sit on top of very different regulatory foundations. Here is the anchor behind each market:

Market	Regulatory stance	Anchor frameworks and key signal
Thailand	Sector decree plus national ethics guidance	1 Royal Decree on Digital Platform Services 2 National AI Ethics Guidelines
The Philippines	Compliance-first, aligned to global standards	1 Data Privacy Act (National Privacy Commission) 2 National AI Strategy Roadmap 2.0 3 EU AI Act alignment
India	Innovation-first with strict data laws	1 India AI Governance Guidelines from the Ministry of Electronics and IT (MeitY), built on seven “sutras” 2 Digital Personal Data Protection Act
China	Hard regulation, gated at market entry	1 Cyberspace Administration (CAC) algorithm registries and security reviews 2 China AI Agent Interconnection standard, June 2026 (unified agent identity)
Hong Kong	Existing-law patchwork, privacy-led	1 Personal Data (Privacy) Ordinance 2 Model Framework from the Privacy Commissioner for Personal Data (PCPD) 3 Generative AI Technical and Application Guideline
Australia	Technology-neutral, existing laws updated	1 Australian Privacy Act automated decision-making amendments (December 10, 2026) 2 Australian AI Safety Institute
Indonesia	Shifting from voluntary to binding	1 Banking AI governance from the financial regulator OJK 2 AI ethics circular (non-binding) 3 Incoming Presidential Regulation on AI
Singapore	Tooling-led market; first in the world with agentic AI guidance	1 Model AI Governance Framework for Agentic AI 2 AI Verify 3 Personal Data Protection Act
Malaysia	Voluntary now, statutory soon	1 AI Governance Bill (pending cabinet) 2 National Guidelines on AI Governance and Ethics 3 MY-AI Standards

What comes next

The organizations that pull ahead will be the ones that can prove what their autonomy did.



For business leaders

Treat accountability as architecture. Decide who owns each AI-driven outcome before deployment, and make Traceability a release requirement. An agent you cannot reconstruct is a liability you cannot price.



For compliance and risk teams

Treat Traceability as an audit requirement. The test is whether a reviewer can reconstruct an AI decision without asking the model to explain itself after the fact.



For regulators and policymakers

The problem is evidence, and the intent is already there. Standards that reward demonstrable Traceability, reconstructable decisions, and tamper-proof trails will move the market faster than rules that only check for written policies. Proof over paperwork.



For the market

Verified accountability is turning into a competitive edge. As agents act and transact at scale, the firms that can show which agent did what, and on whose authority, will be the ones others are willing to transact with. The same proof that satisfies a regulator doubles as a commercial signal: a counterparty is far more willing to let your agent into its systems when that agent carries a verifiable identity and a clean trail. Trust infrastructure binds an action to an authorized agent and a responsible human. Expect it to move from a compliance cost into a reason customers and partners choose you.

Where the market is heading

Three shifts are already in motion:

- 01 Organizations are planning an aggressive move up the maturity ladder.**

52.7% aim to reach a very advanced governance posture, with clear frameworks, strong technical controls, and high confidence, over the next 12 to 24 months, driven by incoming mandates like the EU AI Act, Australia's 2026 automated decision-making rules, and Malaysia's AI Governance Bill.
- 02 Human-to-AI identity traceability is becoming a mandate.**

The direction of travel is clear: 98.6% of organizations rate themselves very or somewhat likely to adopt software that traces AI actions back to a responsible human identity (the identity can be a department or group, as well as a single person). Expect rapid uptake of Identity-as-a-Service and cryptographic logging: if an autonomous agent signs a contract, moves funds, or alters client data, that action gets bound to a verified human signature.
- 03 The model is moving from paper-based checklists to continuous audits.**

As error rates and out-of-bounds operations climb in high-growth hubs like India, the Philippines, and Thailand, boardrooms are accepting that written policies do not stop algorithmic drift. The shift is toward automated logging, real-time guardrails, and third-party certification, all designed to prove compliance dynamically.

AI governance readiness checklist

A practical way to close the Accountability Asymmetry, organized by the **three pillars**. Work down each list before your next agentic deployment.

Autonomy: know what your AI is doing

Map every AI system in production and record the most advanced action each one can take without a human

☐

Document where each system acts for internal teams, for customers, and for external parties

☐

Set a hard boundary against scope creep, so a system cannot quietly drift past its original purpose

☐

Define which actions need human approval before they are performed

☐

Require staged rollout for any new autonomous capability before it runs at full scale

☐

Keep a manual override or pause capability for every autonomous system

☐

Responsibility: own the outcome before it happens

Assign a named owner, a person, department, or group, to every AI-driven outcome before deployment

☐

Document an overview of a given AI system, including its purpose, functioning, and known limitations, and identify the person or team responsible for it

☐

Put an approval gate on high-stakes actions: payments, onboarding decisions, and data changes

☐

Re-check ownership whenever a system's scope changes

☐**Traceability:** prove it after the fact

Maintain a complete list of AI systems in use—a prerequisite for auditing or governing any of them

☐

Make a reconstructable decision pathway a release condition

☐

Keep a tamper-proof audit trail for every autonomous action

☐

Maintain a record of AI-related incidents and how they were resolved

☐

Apply the reviewer test: someone should be able to reconstruct a decision without asking the model to explain itself

☐

Bind each agent action to a verified human identity

☐

Add independent audits or validation on top of internal checklists

☐

How Sumsub can help

Sumsb builds the trust infrastructure for AI that acts. With Summy AI, AI-powered decisioning, MCP connectivity, and pre-built agentic skills at its core, teams can build and adapt compliance workflows faster, from onboarding to ongoing monitoring.



Autonomy: Agent Skills and MCP (Model Context Protocol) integration

Sumsb lets AI agents manage the platform's configuration layer through the Model Context Protocol (MCP). Teams feed changing internal requirements to an agent, which translates them into live, isolated sandbox settings with strict boundaries that maintain human approval for core infrastructure changes. Presently, every agent is supported, including the official connectors for OpenAI ChatGPT and Anthropic Claude.



Responsibility: AI Agent Verification and the KYA framework

Sumsb binds an AI system to a verified, accountable owner (individual or team, whichever matches how the organization assigns responsibility). When an autonomous system starts a high-value action, such as a large transaction or a customer onboarding step, the platform reads the risk level and can trigger an on-demand, targeted Liveness check for an authorized person acting on behalf of that owner. Every automated action then traces back to a real, accountable owner.



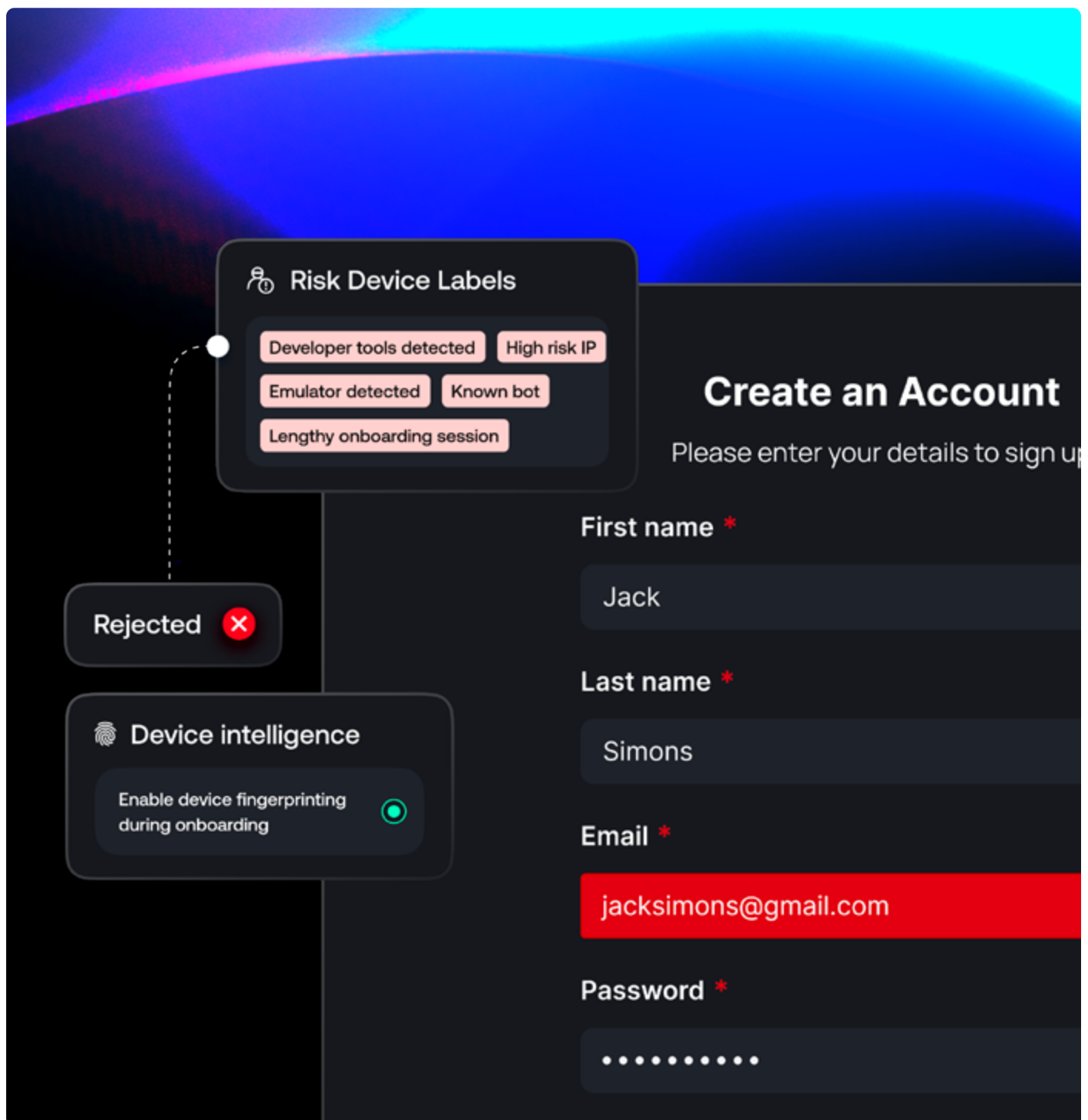
Traceability: Summy, the AI Copilot

Grounded in platform data, Summy is Sumsb's platform-native AI investigator. It replaces manual tracking by summarizing complex case anomalies into clear, audit-ready pathways, so every automated flag stays traceable and explainable when a regulator asks.

Supporting controls

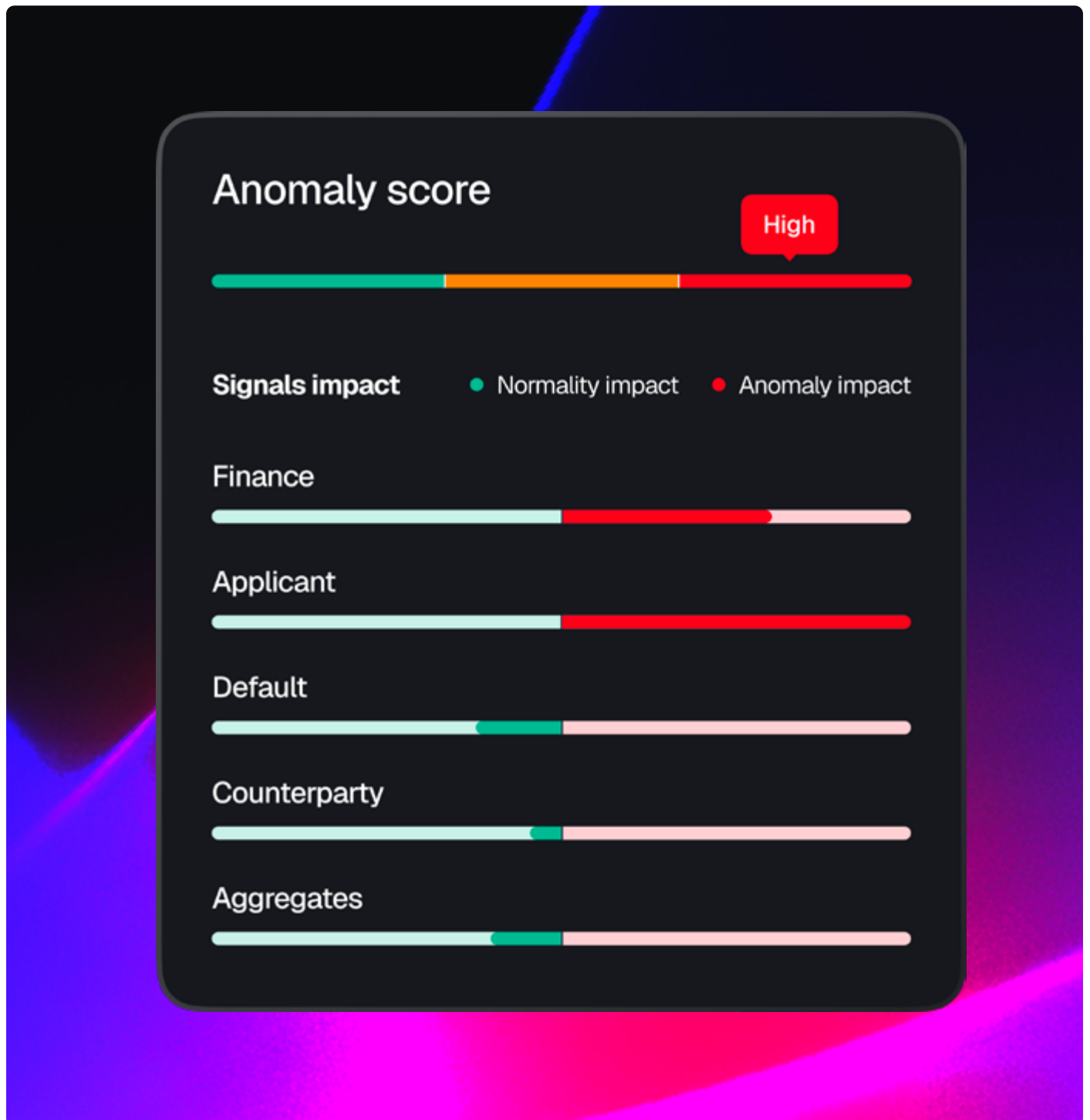
Device Intelligence and bot detection

Sumsb tracks device, session, and behavioral signals in real time to intercept rogue automated processes and API misconfigurations before a data pipeline is contaminated or compromised.



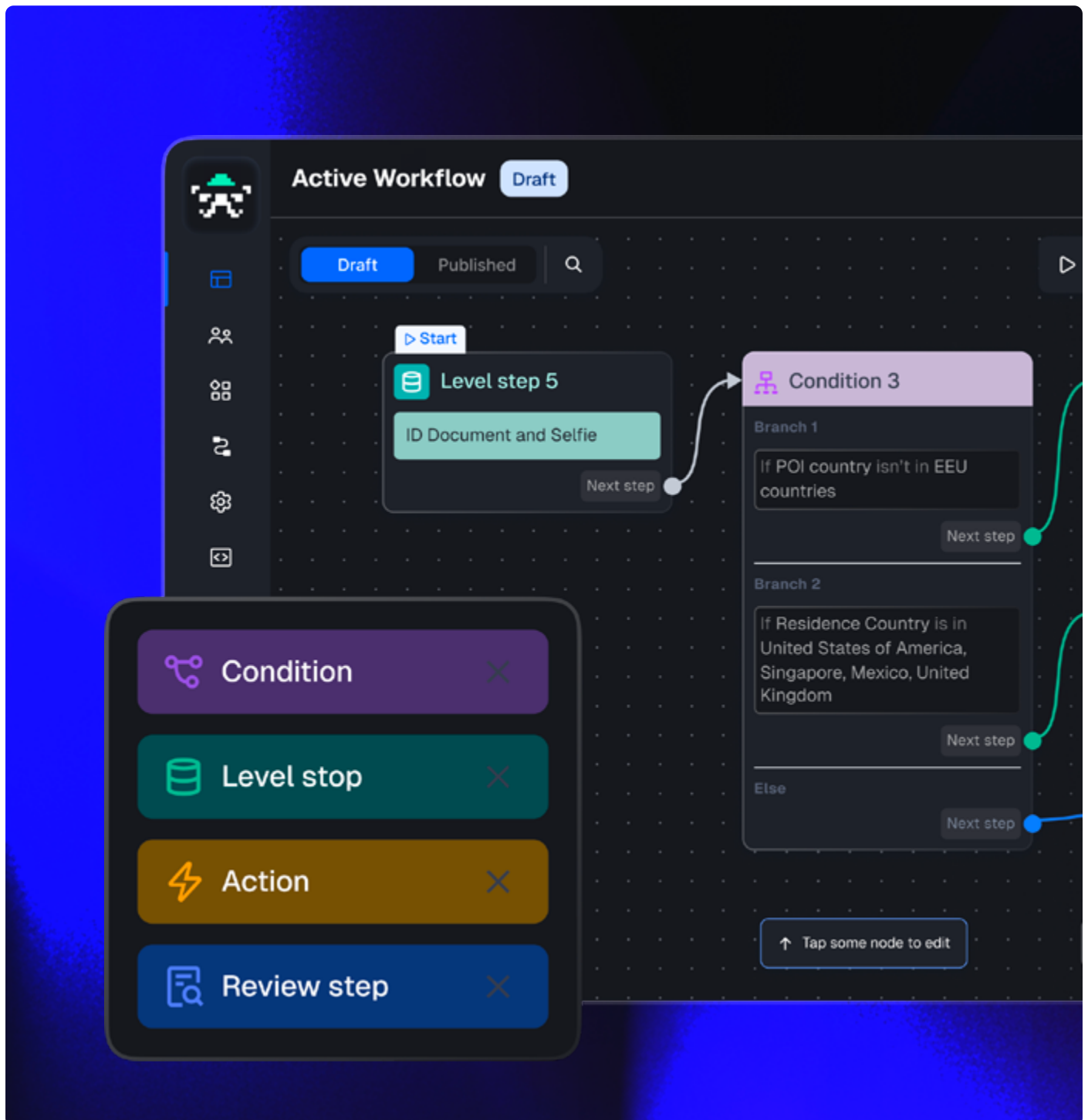
Money Muling Prevention and Risk Scoring

Fraudsters run clusters of automated browser profiles to mimic real users, so basic IP blocks fail. Sumsu reads behavioral signals across accounts and sessions to expose coordinated mule networks and the bots behind fake shoppers and promo abuse.



Workflow Orchestration

Compliance teams build fraud rules through plain-language queries instead of engineering tickets, tuning them to local risk levels, so lower-margin businesses get fraud coverage matched to their market without a dedicated, high-cost data-engineering team.



Keep your business safe from advanced fraud with Sumsub



User Verification



Business Verification



Transaction Monitoring



Fraud Prevention



Device Intelligence

95 %

AML matches resolved
with AI assistance

690 k+

fraud attempts
prevented monthly

<59 sec

to verify a user



See Sumsub in action

[Book a free demo](#)