

FINANCIAL SERVICES · INSURANCE, WEALTH AND BANKING

Partial Understanding

Artificial intelligence across regulated financial services, and the accountability that has not yet been allocated

Paul Forrest

Fractional Head of AI, Ecaveo
September 2026

About this paper

This is the financial services paper in the Ecaveo Responsible AI Series. It sits alongside a trilogy on professional firms, covering private equity, accounting and law, and a paper on freight and warehouse operations. Each stands alone.

The title is taken from a survey category. When the Bank of England and the Financial Conduct Authority asked 118 regulated firms in 2024 how well they understood the artificial intelligence they were running, 46 per cent selected partial understanding and 34 per cent selected complete.⁷ **REGULATOR** That answer, given by the firms themselves to their own supervisors, is the honest starting point for everything in this paper.

Editorial balance. Roughly seven parts in ten of this document is evidence and sector analysis, two parts in ten is original framework, and one part in ten concerns how I would run the work. Where the evidence is thin the paper says so, and in this sector it is thinner than the volume of commentary suggests.

HOW TO READ THE EVIDENCE TAGS

Every substantive claim carries a tag stating what kind of evidence supports it.

REGULATOR A regulator, government department or official statistics publisher speaking in its own voice.

STANDARD A published standard, statute, statutory instrument or formal guidance document.

CASE LAW A judgment, tribunal decision or enforcement action.

PEER-REVIEWED A primary study, systematic review or parliamentary investigation working from primary material.

SURVEY A survey or dataset with disclosed method and sample.

VENDOR RESEARCH Research from a body with a commercial or membership interest in the answer. Directional only.

PRACTICE Practitioner opinion, including my own.

EVIDENCE GAP A question the available evidence does not answer.

ECAVEO My own framework, proposed rather than proven.

Legal and regulatory notice. This paper describes legislation, regulatory rules, supervisory guidance and enforcement action throughout, and Chapters 2, 3 and 4 in particular. Nothing in this document constitutes legal advice, and the provision of legal advice sits outside the terms of any engagement with the author or with Ecaveo. The material is presented to support discussion and further review by qualified advisers.

Statements about what a source says are attributed to that source. Where I draw a practical implication the source does not itself state, the text says so. The regulatory position described here moved several times during 2026 and is likely to move again.

Contents

	A note from the author	4
	Executive summary	5
ONE	Adoption without comprehension	9
TWO	What the regulator has declined to do	13
THREE	The customer outcome test	16
FOUR	Meaningful human involvement	20
FIVE	Where the value is credible	23
SIX	The model you did not know you had	26
SEVEN	Three deployments, and how to evaluate them	29
EIGHT	The economics of verification	32
NINE	The first 90 days, then the rest of the year	35
TEN	Where the evidence is thin	37
	References	39

THE ARGUMENT IN ONE PARAGRAPH

Three quarters of surveyed UK financial firms are using artificial intelligence, and almost half tell their supervisors they understand it only partially. The Financial Conduct Authority has decided that no new rules are needed, resting instead on the Consumer Duty and the Senior Managers and Certification Regime. Yet no senior management function and no prescribed responsibility for artificial intelligence has been created, so the accountability the strategy depends on has been asserted and not allocated. This paper argues that the allocation is the work. A firm that names the accountable individual, writes down what that person is accountable for, and builds the evidence that lets them discharge it, will find every subsequent rule cheap. A firm that waits for the rule will be building the evidence under deadline.

ILLUSTRATIVE MATERIAL

Every firm, adviser, claim and customer in the worked examples is invented, and each is marked as an illustrative scenario where it appears. Published research, official statistics and enforcement decisions are cited and numbered. Invented material is never presented as evidence.

CITATION

Forrest, P. (2026). *Partial Understanding: Artificial Intelligence in Insurance, Wealth Management and Financial Services*. Ecaveo Responsible AI Series, Financial Services. Version 1.0, September 2026.

ALSO IN THIS SERIES

The Traceable Thesis. Private equity.

Sufficient and Appropriate. Accounting and audit.

Verified Authority. Law firms.

The Protected Constraint. Operations.

A NOTE FROM THE AUTHOR

Why this paper is about accountability and not about models

The interesting question in financial services is no longer whether the technology works. It is who answers for it, and that question has an unsatisfying answer at the moment.

My background is business process work from the 1990s, then machine learning and language models from 2014. I build small specific models in preference to large general ones, on the view that a model trained on everything is calibrated for nothing. Most of my work is with private equity and venture capital portfolio businesses, which means I see financial services from two directions at once, as a regulated sector and as a set of companies somebody owns and expects a return from.

What prompted this paper was a mismatch. Read the FCA's own material and you find a regulator that has thought carefully, engaged widely and reached a defensible conclusion, which is that existing frameworks are adequate and new AI rules would be premature.^{5,6} **REGULATOR** Read the Treasury Committee's report from January 2026 and you find the same approach described as one that is exposing consumers and the financial system to potentially serious harm.¹⁸ **PEER-REVIEWED** Both of those documents are careful. They are looking at the same sector and reaching opposite conclusions about the same policy.

The disagreement is not really about technology. It turns on whether an accountability regime designed for human decisions transfers cleanly to systems that make decisions faster than anybody can review them, and whether saying that senior managers are accountable is the same thing as making them accountable. My own view, which the rest of this paper tries to earn, is that it is not the same thing, and that the gap between the two is where the work sits.

I have tried throughout to separate three kinds of statement, which are what a source says, what I infer from it, and what I would do. The evidence tags do most of that work. Where I have gone beyond a source I have said so in the sentence rather than in a footnote.

TWO PAGES THAT CIRCULATE ON THEIR OWN

Executive summary

Adoption is mainstream, comprehension is not, and the accountability regime the strategy relies on has a hole in the middle of it.

The position in September 2026 can be stated without much hedging. Firms have adopted. Supervisors have decided against new rules. Parliament has objected. Nobody has measured whether any of it works, and the one rigorous UK evaluation of an algorithmic intervention in this sector found a null result in half the market it examined.

1

Three quarters of firms use it and almost half say they understand it partially. The joint Bank of England and FCA survey of 118 regulated firms, published in November 2024, found 75 per cent already using artificial intelligence and a further 10 per cent planning to within three years. On understanding, 46 per cent reported partial understanding of the technologies they had deployed against 34 per cent claiming complete understanding. A third of live use cases were third-party implementations.⁷ **REGULATOR** No later edition of this survey exists, which means the sector's own supervisors are working from a picture that is now almost two years old.

2

The FCA has ruled out AI-specific rules, in terms. In a blog of 8 June 2026 the FCA wrote that it has “been clear that we are not going to introduce new regulations for AI” and will instead rely on existing frameworks including the Consumer Duty and the Senior Managers and Certification Regime.⁶ **REGULATOR** The Financial Services AI Adoption Plan published by HM Treasury on 14 July 2026 recommends that regulators confirm the same principles-based, outcomes-focused approach, which is a recommendation and not yet a confirmation.¹⁹ **ADVISORY**

3

The accountability vehicle has no AI seat in it. No senior management function and no prescribed responsibility for artificial intelligence has been created under the Senior Managers and Certification Regime. The Treasury Committee asked the FCA to clarify senior manager accountability for AI by the end of 2026. In April 2026 the FCA declined to commit to comprehensive guidance, and as at the date of this paper no clarification has been published.^{5,18,20} **REGULATOR**

4

Parliament has said the approach is unsafe. The Treasury Committee's report of 20 January 2026 concluded that the wait-and-see approach taken by the FCA, the Bank of England and HM Treasury is “exposing consumers and the financial system to potentially serious harm”.¹⁸ **PEER-REVIEWED** It recommended AI-specific stress testing and the designation of major AI and cloud providers as critical third parties by the end of 2026.

5

Only the cloud half of the concentration problem has been addressed. HM Treasury designated four critical third parties with effect from 13 July 2026, all of them cloud and infrastructure providers.²¹ **REGULATOR** No model provider has been designated. The 2024 survey found the top three model providers accounted for 44 per cent of models named by firms, so the concentration the Committee was worried about sits largely outside the new regime.⁷ **REGULATOR**

6

No independent measurement of return exists for this sector. The Office for National Statistics excludes finance and insurance from the business survey that produces the UK's official AI adoption statistics, so there is no independent official figure at all.²² **REGULATOR** I could find no randomised or quasi-experimental study measuring AI effects in underwriting, claims handling or regulated advice. **EVIDENCE GAP**

What follows for a board

Name the person. Under the current regime there is no prescribed responsibility to point at, so the allocation has to be made internally and written down. That means a named senior manager whose statement of responsibilities says, in words, that they are accountable for the firm's use of artificial intelligence in regulated activity. Doing this now costs a paragraph. Doing it after the FCA publishes its expectations costs a programme.

Build the inventory before the validation. A third of use cases are somebody else's implementation and an unknown further number are features switched on inside software the firm already licenses. Nobody can validate a model they have not found. The discovery process matters more, at this stage, than the validation methodology.

Test outcomes by cohort, and keep the denominator. The FCA's own evaluation of its general insurance pricing remedies could not say which customer groups benefited, because the data to answer that question was never collected.¹⁶ **REGULATOR** A firm that designs cohort measurement into a deployment on day one can answer the question a supervisor will eventually ask. A firm that does not, cannot answer it retrospectively at any price.

Chapter 9 sets out a 90-day proof across three bounded workflows, run in shadow mode. Chapter 10 states what would change the position taken here.



PART ONE · THE POSITION

**Adoption has outrun comprehension,
and measurement has not started**

CHAPTER ONE

Adoption without comprehension

The most quotable governance finding in UK financial services was volunteered by the firms themselves, to their own supervisors, and it has not been updated since.

In November 2024 the Bank of England and the Financial Conduct Authority published the third of their joint surveys of artificial intelligence in UK financial services. The instrument went to regulated firms across six sectors, covering UK banks, international banks, insurers, non-bank lenders, investment and capital markets firms, and financial market infrastructure and payments. It drew 118 responses.⁷ **REGULATOR** Being a regulator survey of regulated firms, it carries a credibility that no vendor study in this sector can match, and it has one methodological weakness that should be stated at the outset, which is that the response rate is not published, so self-selection cannot be excluded and the headline percentages count firms without weighting them by size.

The headline is familiar. Adoption stood at 75 per cent, with a further 10 per cent planning to adopt within three years, up from 58 per cent using and 14 per cent planning in 2022.⁷ **REGULATOR** That number has been quoted in almost every subsequent official document, including HM Treasury's adoption plan of July 2026 and the Treasury Committee's report of January 2026.^{18,19} **REGULATOR**

The finding underneath it has been quoted much less. Asked about their understanding of the AI technologies they use, 46 per cent of firms reported partial understanding. Complete understanding was claimed by 34 per cent.⁷ **REGULATOR** Set those two figures next to each other and the picture changes. This is not a sector experimenting cautiously at the edges. It is a sector that has deployed at scale and, on its own account, does not fully understand a substantial part of what it has deployed.

Materiality softens that a little. The same survey found 62 per cent of use cases rated low materiality and 16 per cent rated high.⁷ **REGULATOR** Most of what firms are running is unremarkable, and partial understanding of an unremarkable system is an ordinary operational condition. The problem is that partial understanding does not distribute itself neatly by materiality, and the survey does not cross-tabulate the two. Whether the 46 per cent and the 16 per cent overlap is exactly the question a supervisor would want answered, and the published data cannot answer it. **EVIDENCE GAP**

THE CORE CLAIM

Comprehension has not kept pace with deployment, and the firms know it. Everything else in this paper follows from that admission.

1.1 The number nobody has refreshed

There is no 2025 or 2026 edition of this survey. The Treasury Committee in January 2026 and the adoption plan published by HM Treasury in July 2026 were both still citing the November 2024 figures as the current picture.^{18,19} **REGULATOR** The adoption plan also records that further survey work is under way at the FCA and the Bank to provide an updated view. For a technology that the same documents describe as moving quickly, the UK's principal supervisory dataset is now approaching two years old, and it predates the general availability of the agentic systems that both IOSCO and the Financial Stability Board have since identified as a distinct category of risk.^{12,26} **REGULATOR**

Alternatives are worse. The Office for National Statistics runs the Business Insights and Conditions Survey, which is the source of the UK's official AI adoption statistics and which reported roughly 35 per cent adoption among businesses with 10 or more employees in its wave of June 2026, from 38,637 responses at a response rate of 26.7 per cent.²² **REGULATOR** Finance and insurance are explicitly outside the survey's coverage. So the only genuinely independent official statistic on AI adoption in the UK does not cover this sector, and no independent official figure for financial services exists.

That absence is not a small technical point. Every adoption figure quoted about UK financial services is either a regulator's own survey of the firms it supervises, a trade body counting its members, or vendor research. Each has a reason to answer in a particular direction, and only the first has a supervisory obligation to be accurate.

1.2 Two adoption figures, both correct

The adoption plan cites two numbers within a few pages of each other. One is the 75 per cent from the joint survey. The other is 21 per cent, from a Department for Science, Innovation and Technology survey in early 2025, for firms in the financial and real estate sectors.¹⁹ **ADVISORY** Presented alone, either misleads. The joint survey samples large regulated firms; the DSIT survey samples businesses generally, and its sector grouping folds financial services in with estate agency.

Reading them together produces the more useful statement, which is that adoption in this sector is concentrated in large firms and thin in small ones. The FCA's own wealth management data supports that reading. Its survey of roughly 400 wealth management firms found 13 per cent using in-house or third-party AI tools, rising to 45 per cent once firms considering it in the following 12 months were included.¹⁰ **REGULATOR** Note the reference date, though. Those figures describe the position as at 31 December 2024 and were published on 18 August 2026, a gap of nearly 20 months.

THE ADOPTION EVIDENCE, GRADED

SOURCE	FIGURE	WHAT IT IS	WEAKNESS
Bank of England and FCA, Nov 2024	75% using, 10% planning	Regulator survey of 118 regulated firms across six sectors	Response rate unpublished; firm count not weighted by size; 22 months old
Office for National Statistics, Jul 2026	c.35% of businesses with 10+ staff	Official statistics, 38,637 responses, 26.7% response rate	Finance and insurance are outside the survey's coverage
DSIT, early 2025, via the adoption plan	21% of financial and real estate firms	Government survey, reported second hand	Sector grouping conflates finance with real estate
FCA wealth management survey, Aug 2026	13% using, 45% using or considering	Regulatory data collection from c.400 firms	Data as at 31 December 2024, published 20 months later
IOSCO consultation report, Mar 2025	c.50% investing or adopting, c.41% in production	Industry survey of IOSCO affiliate members, published alongside a separate survey of 24 of 33 regulators	Self-reported by firms; global rather than UK

TABLE 1 Every published adoption figure for UK financial services, with its provenance. The absence of any independent official statistic is the finding, and it follows from the ONS excluding finance and insurance from its business survey.

1.3 What the technology is measured to do, and where

Away from the sector, the evidence on what these systems do to work is good and it points in two directions at once. That divergence is the most useful thing a board can be told, because it is stable across studies and it predicts which deployments will disappoint.

On the positive side, a staggered rollout of a generative AI assistant across 5,179 customer support agents produced a 14 per cent average increase in issues resolved per hour, concentrated heavily among novice and low-skilled workers at 34 per cent, with minimal effect on experienced high performers.²⁷ **PEER-REVIEWED** Three randomised field experiments with 4,867 software developers at Microsoft, Accenture and a third large firm found a 26 per cent increase in completed tasks, again with larger gains for less experienced developers.²⁸ **PEER-REVIEWED** The second study should be read with its conflicts visible, since the tool studied is a Microsoft product and two of the three sites are the vendor and a large systems integrator.

Against that sits a randomised experiment with 758 consultants, which found a 12.2 per cent increase in tasks completed and 25.1 per cent faster completion on tasks inside the model's competence, and participants 19 per cent less likely to produce a correct answer on tasks outside it.²⁹ **PEER-REVIEWED** And a small randomised trial of 16 experienced open-source developers across 246 real tasks found that allowing AI increased completion time by 19 per cent.³⁰ **PEER-REVIEWED** That study has severe limitations, with a tiny sample of self-selected experts working on codebases they knew deeply, and it should not be generalised. One finding inside it travels anyway. The developers forecast a 24 per cent speedup beforehand, and after the experiment still believed they had been 20 per cent faster while the data showed them 19 per cent slower.

Hold that last point next to how AI pilots are actually evaluated in financial services, which is overwhelmingly by user satisfaction survey. A 19 per cent slowdown experienced as a 20 per cent speedup is not a curiosity. It is a direct challenge to the measurement method on which most business cases in this sector currently rest.

1.4 The gap in the sector's own evidence

Here is the finding I expected to be able to avoid writing. I could locate no randomised or quasi-experimental study measuring the effects of artificial intelligence in insurance underwriting, claims handling, or regulated financial advice. **EVIDENCE GAP** The machine learning literature on anti-money laundering and fraud consists of model-performance papers evaluated retrospectively on historical datasets. Those papers typically report reductions in false positives without measuring whether real criminal activity was detected more often, or whether legitimate customers were harmed less.

So anybody asserting a measured productivity or accuracy return from AI in UK underwriting or claims is extrapolating from adjacent domains. That is a legitimate thing to do, provided it is said out loud. The customer support and consulting studies above are the best available proxies, and they suggest a return that is real, concentrated among less experienced staff, and reversed on tasks that fall outside the system's competence. Which tasks those are, in an insurance claims workflow, nobody has published.

CHAPTER TWO

What the regulator has declined to do

The FCA's position is coherent, stated plainly, and contested by Parliament. Understanding the disagreement matters more to a board than picking a side in it.

British financial regulation has taken a decision about artificial intelligence, and the decision is to make no new rules. This is not drift or delay. It is a considered position that the FCA has stated repeatedly and defended when challenged.

The clearest statement came in a blog by the FCA's Head of Cross-cutting Policy and Strategy on 8 June 2026. "We have been clear that we are not going to introduce new regulations for AI. Instead, we'll rely on existing frameworks, including the Consumer Duty, the Senior Managers and Certification Regime (SM&CR), and our expectations on governance and controls."⁶ **REGULATOR** The FCA's standing approach page, last updated 13 February 2026, gives the reasoning, which is that with a fast-moving technology this is the best way to support UK growth and competitiveness.⁵ **REGULATOR**

Two things follow that are worth being precise about. The FCA's approach page contains no numbered set of AI principles and no framework taxonomy, so nothing of that kind should be attributed to it. And the Consumer Duty consultation actually running at the time of writing, CP26/23 of 29 June 2026 on scope and proportionality, contains no AI-specific provisions at all.⁹ **REGULATOR** The link between AI and the Consumer Duty is made in the FCA's commentary. It has not been made in the rules.

2.1 The objection from Parliament

On 20 January 2026 the Treasury Committee published its report on artificial intelligence in financial services. Its central finding is unusually direct for a select committee. The three authorities are, in the Committee's words, "exposing consumers and the financial system to potentially serious harm" through a "wait-and-see approach".¹⁸ **PEER-REVIEWED**

The Committee identified opacity in AI-driven credit and insurance decisions, financial exclusion, unregulated AI financial advice, increased fraud, AI-enabled cyber attack, dependence on a small number of United States technology providers, and AI-driven trading amplifying herding behaviour. It made three recommendations. The FCA should publish comprehensive guidance by the end of 2026 on how consumer protection rules apply to AI and on senior manager accountability. The Bank and the FCA should conduct AI-specific stress testing. HM Treasury should designate major AI and cloud providers as critical third parties by the end of 2026.¹⁸ **PEER-REVIEWED**

The responses came on 16 April 2026 and they are instructive.²⁰ **REGULATOR** The Bank of England accepted the stress testing recommendation in principle, describing scenario analysis under way and work to incorporate AI scenarios into cyber and operational testing. HM Treasury deferred on designation, saying it was gathering evidence and expected initial decisions during the year, and then delivered in part. The FCA partially accepted, declined comprehensive new guidance, and cited its principles-based and outcomes-focused approach along with what it described as widespread industry support for using the existing framework.

Read that last response carefully. The regulator's argument for not writing AI guidance is, in part, that the regulated population does not want it. That may well be true. It is also not

usually treated as dispositive when the question is consumer protection, and the Committee's point was precisely that the people who would be harmed are not the people being surveyed.

2.2 The seat that was never created

The Senior Managers and Certification Regime is the mechanism on which the whole no-new-rules strategy depends. Its logic is that firms allocate responsibilities to named individuals, those individuals hold a statement of responsibilities, and the regulator can then hold a person rather than an institution to account when something goes wrong.

No senior management function for artificial intelligence exists. No prescribed responsibility for artificial intelligence exists. The FCA has not created either, and its own approach page goes no further than saying that its rules emphasise accountability for senior managers and are relevant to the safe use of AI.⁵ **REGULATOR** The Treasury Committee asked for clarification by the end of 2026 and the FCA declined to commit to it. As at the date of this paper no clarification has appeared.^{18,20} **REGULATOR**

My own reading, which goes beyond what either document states, is that this is the load-bearing weakness in the current settlement. An accountability regime works by allocation. Where a responsibility has not been allocated, the regime does not produce accountability, it produces an argument after the event about which of several senior managers should have owned the thing. I suspect that argument, when it comes, will be resolved against whichever individual holds the operational function that failed, and that this is a worse outcome for firms than a prescribed responsibility would have been.

2.3 What the FCA is doing instead

The absence of rules is not an absence of activity, and a board should know what the regulator is actually building, because it shapes what supervisory conversations will look like in 2027.

AI Live Testing runs firms' systems in supervised conditions. The engagement paper appeared on 29 April 2025 and the feedback statement, drawing on 67 responses, on 9 September 2025. A second cohort of eight firms was announced on 21 April 2026, including Barclays, Experian, Lloyds Banking Group through Scottish Widows, and UBS, working on targeted support for investments, credit score insights, agentic payments, anti-money laundering detection and know-your-customer processes. The first cohort included NatWest, Monzo and Santander.²³ **REGULATOR** The evaluation report is due at the end of 2026 or early 2027, so no outcomes have been published. Alongside it the FCA has twice committed to publishing examples of good and poor practice on AI during 2026, in the June blog and again in the Mills Review press release, and its AI Input Zone call for views closed on 19 June 2026.^{6,17} **REGULATOR**

The Mills Review, published on 6 July 2026 and led by Sheldon Mills, is the more consequential document. It commissioned research from Yonder Consulting in April 2026 covering more than 5,000 UK retail financial services consumers, quota-controlled for age, gender, ethnicity, region, housing tenure and internet ability.¹⁷ **REGULATOR** Its headline is that 20 per cent of UK adults, roughly 11 million people, are likely to use autonomous AI within pre-set financial goals.

Be careful with that figure, because the same review is rendered two ways across the FCA's own documents. The wealth management report gives it as one in five UK adults being "already open to AI making financial decisions for them". The press release for the Mills Review frames it as being likely to use autonomous AI within pre-set financial goals.^{10,17} **REGULATOR** Openness to delegation and likelihood of adoption within stated constraints are different propositions, and quoting whichever suits the argument is a temptation

WORTH HIGHLIGHTING

The FCA's own page claims only that its rules "emphasise accountability for senior managers and are relevant to the safe use of AI". That is a statement about relevance, not about allocation.

the source itself makes easy. Both describe intention. Survey research always does, and it is worth remembering when the figure appears on a strategy slide.

THE REGULATORY POSITION IN SEPTEMBER 2026

01 Stated No new AI rules. Existing frameworks apply, principally the Consumer Duty and the Senior Managers and Certification Regime.^{5,6}	02 Contested The Treasury Committee finds the approach is exposing consumers and the system to potentially serious harm.¹⁸	03 Partly delivered Four cloud providers designated as critical third parties from 13 July 2026. No model provider designated.²¹	04 Outstanding Senior manager accountability for AI remains unclarified. Good and poor practice examples are promised for 2026 and have not appeared.^{18,20}
---	---	---	---

FIGURE 1 The four positions a board should be able to state.

Compiled from FCA publications, the Treasury Committee's fifteenth report of session 2024 to 2026, the responses published on 16 April 2026, and the Bank of England announcement of 10 July 2026.

2.4 The European question, and a date that has moved

The EU AI Act reaches UK firms through Article 2. It applies to providers placing an AI system on the Union market irrespective of where they are established, and to providers and deployers located in a third country where the output produced by the system is used in the Union.² **STANDARD** A UK insurer or wealth manager is therefore in scope if its AI output is consumed in the Union, and establishment in the United Kingdom does not answer the question.

Annex III captures two financial services categories, and the carve-outs matter. Point 5(b) covers systems used to evaluate creditworthiness or establish a credit score, with fraud detection expressly excluded. Point 5(c) covers risk assessment and pricing in relation to natural persons in the case of life and health insurance.² **STANDARD** General insurance pricing for motor or home cover is therefore not caught by that limb. This distinction is routinely lost in commentary and a specialist reader will notice if it is.

A date many boards were briefed on has changed. Regulation (EU) 2026/1744, the Digital Omnibus on AI, was adopted on 8 July 2026, published in the Official Journal on 24 July and entered into force on 27 July 2026. It defers the obligations for stand-alone Annex III high-risk systems to 2 December 2027, and for high-risk systems embedded in regulated products to 2 August 2028. These are fixed dates with no conditionality.² **STANDARD** A firm still planning against August 2026 is planning against a repealed timetable. The prohibitions, the AI literacy duty, the general-purpose model rules and the transparency obligations remain in force now.

The Omnibus made several other changes worth a line. The AI literacy obligation was softened from providers ensuring staff literacy to supporting its development. The definition of a safety component was narrowed. Delegated acts may limit the Act's application where sectoral law is equivalent, which is the provision most likely to matter to financial services over the next two years.² **STANDARD**

CHAPTER THREE

The customer outcome test

The Consumer Duty is an outcomes regime, which means a firm has to show what happened to customers. That is a data problem before it is a compliance problem, and most firms have not solved it.

The Consumer Duty asks firms to design products and services that meet the needs of their target customers and provide fair value, to communicate in a way that meets customers' information needs, and to provide support that meets customers' needs.⁵ **REGULATOR** Nothing in that formulation mentions technology, which is exactly why the FCA regards it as adequate for AI. An outcome is an outcome whatever produced it.

The difficulty is evidential. To show that an outcome was met, a firm needs to know which customers received which treatment and what happened to them afterwards, broken down by the characteristics that matter. Most AI deployments in this sector have not been instrumented to produce that data, and it cannot be reconstructed later.

3.1 The best-evidenced algorithmic harm in the UK, and its remedy

General insurance pricing is the one place where the UK has a full cycle of evidence, running from quantified harm through intervention to independent evaluation. It is worth walking through, because the shape of the evidence is more instructive than the conclusion.

In May 2021 the FCA published PS21/5, its final rules following the general insurance pricing practices market study. The harm was quantified as follows. Six million policyholders paid high prices in 2018, and if they had paid the average for their risk they would have saved £1.2bn.¹⁵ **REGULATOR** The remedy was a pricing rule requiring that the renewal price offered to an existing customer should be no greater than the equivalent new business price. It came into force for the pricing remedy on 1 January 2022.

The FCA then did something regulators rarely do, which is evaluate its own intervention properly. EP25/2, published in July 2025, uses a continuous difference-in-differences design with treatment intensity defined by the degree of price-walking before the rule, measured at the level of a policy grouping. It covers 16 home insurers at roughly 80 per cent market share and 13 motor insurers at roughly 57 per cent, from the first quarter of 2019 to the first quarter of 2024, and the parallel trends assumption held in the pre-intervention period.¹⁶ **REGULATOR**

In home insurance, the gap between the price paid by existing and new customers almost halved, from £95.38 to £49.17. In motor the gap ran the other way to begin with and then widened, since new customers who had previously paid £20.76 more than existing ones now pay £109.19 more, which the FCA reads as risk-appropriate pricing reasserting itself. Motor prices fell by a statistically significant average of £6.63 per policy, giving a central ten-year monetised saving of £1.6bn against a range from £163m to £3.0bn. In home insurance there was no statistically significant effect on prices at all.¹⁶ **REGULATOR**

Two features of that evaluation deserve more attention than the headline. The first is the null result. A major intervention in a market where substantial harm had been quantified produced no measurable price effect in half the market examined, and the FCA reports it plainly. The second is the limitation the FCA states in its own words, which is that due to data limitations it is unable to specify which customer groups have benefited more and which less.¹⁶ **REGULATOR**

WHY THIS CASE MATTERS

It is the only UK intervention in algorithmic pricing that has been evaluated with a credible causal design. The results are mixed, and the mixture is the lesson.

Sit with that for a moment. The regulator intervened to protect customers who were being penalised for loyalty, on the reasonable assumption that those customers skew towards the inattentive, the elderly and the vulnerable. Four years later it cannot say whether they were the ones who benefited. The analysis was sound. Nobody had collected the data that would answer the question. Any firm deploying a pricing or eligibility model today is one intervention away from being asked exactly this, and the answer will depend on decisions made at deployment.

3.2 Disparate outcomes, and the limits of what the evidence shows

Citizens Advice published research in March 2022 reporting that people of colour paid on average £250 a year more, across 18,000 car insurance costs drawn from its debt clients in 2021, and at least £280 a year more in postcodes where more than half the population were people of colour. It also commissioned 649 mystery-shopping exercises across eight postcodes using six personas.³¹ **VENDOR RESEARCH**

That is a serious piece of work and it should be handled carefully. It is charity-commissioned observational research, the postcode analysis cannot separate ethnicity from correlated geographic risk factors, and the press release records no FCA response. It is legitimate evidence of a disparate outcome and of a question the regulator has not answered. It is not evidence of unlawful discrimination and should never be cited as though it were.

The academic literature is more precise and less comfortable. Two studies from the United States mortgage market are the strongest causal work available. The first found that lenders charge Latinx and African-American borrowers 7.9 basis points more on purchase mortgages and 3.6 basis points more on refinancing, costing those borrowers around \$765m a year, and that fintech algorithms discriminate 40 per cent less than face-to-face lenders while not discriminating at all in the approve or decline decision.³² **PEER-REVIEWED** The second, a counterfactual simulation using US mortgage data, found that Black and Hispanic borrowers are disproportionately less likely to gain from the introduction of machine learning, and that the increase in disparity is primarily attributable to greater model flexibility.³³ **PEER-REVIEWED**

The mechanism in the second study is the part that matters for a UK insurer. The disparity does not arise from the model triangulating protected characteristics. It arises from flexibility, because a more flexible model can find and price fine-grained risk distinctions that a simpler model averaged over. Removing protected variables does not address that. Neither does removing obvious proxies. The only thing that detects it is outcome measurement by cohort after deployment, which returns us to the data problem in section 3.1.

Both studies are US mortgage market work from the pre-LLM era, and neither transfers automatically to UK insurance or advice. I include them because they are the best-identified evidence in existence on this question and because the flexibility finding is a mechanism rather than a correlation, which makes it more portable than most.



Underwriting and claims. The workflows where a model's output most directly determines what a customer receives, and where cohort measurement is hardest to retrofit.

3.3 Where the complaints are, and where they are not

The Financial Ombudsman Service recorded 214,600 new complaints in 2025/26, down almost 30 per cent from 305,700 the year before, with an uphold rate of 30 per cent.²⁴ **REGULATOR** Resist the temptation to read that fall as improving conduct. It is largely an artefact of a policy change, since professional representatives fell from 50 per cent of cases to 18 per cent after the FOS introduced referral charges.

More useful for our purposes is what the FOS says about AI directly. In its response to the Mills Review in February 2026 the service wrote that “at present, we are receiving very few complaints about a firm’s use of AI”, while acknowledging that around three quarters of UK firms are now deploying it.²⁵ **REGULATOR** The service also flags that rigid or opaque algorithms could risk filtering out complex or non-standard vulnerabilities.

That near-absence cuts two ways and the FOS does not claim to know which. Either AI harm is not yet materialising in a form that produces complaints, or it is materialising and is not being detected or attributed. The published complaints dataset contains no AI or automated-decision category, so there is currently no way to measure AI-related complaint volumes from FOS data at all. **EVIDENCE GAP**

The FOS is clearer about what it expects. Consumers “must be able to understand how outcomes were reached and receive clear explanations in plain English”, and transparency, fairness and the ability to challenge an outcome are described as foundational.²⁵ **REGULATOR** A firm that cannot produce a plain-English explanation of an adverse automated outcome, on request, six months later, has a problem with the ombudsman regardless of what the FCA has or has not written.

3.4 The other direction of travel

One of the odder features of 2026 is that generative AI has arrived in the complaints system from the consumer side faster than from the firm side. In a small sample review the FOS found that up to 35 per cent of responses to its initial assessments appeared AI-generated or heavily AI-assisted. It documented hallucinations including fabricated laws and invented decisions, and submissions running beyond 200 pages in response to a provisional decision of six to eight pages.²⁵ **REGULATOR** The service is explicit that this is early small-sample analysis and the 35 per cent is an estimate.

I mention it partly because it is the best-documented instance of generative AI causing measurable harm inside the UK financial services system, and partly because it changes the operational picture for complaint handling teams in a way that no strategy document has yet acknowledged. The volume is arriving now. Whether firms should respond by using the same tools is a question I will come back to in Chapter 8, and I am not sure the answer is yes.

CHAPTER FOUR

Meaningful human involvement

The statutory test for whether a decision counts as automated changed in February 2026, and the phrase that now governs it has not been defined by anyone.

Section 80 of the Data (Use and Access) Act 2025 replaced Article 22 of the UK GDPR with a new Section 4A containing Articles 22A to 22D. It came into force on 5 February 2026.⁸ **STANDARD** Anybody briefing a board that automated decision-making rules are unchanged is working from pre-February material.

The change is real and it is narrower than the headlines suggested. Under the old Article 22, a solely automated decision with legal or similarly significant effect was prohibited unless one of three gateways applied, being contract, explicit consent, or authorisation by law. Under Article 22A and 22B, most of the Article 6 lawful bases become available for such decisions where special category data is not involved, with recognised legitimate interests under Article 6(1)(ea) expressly carved out by Article 22B(4). The ICO's own summary is that the Act opens up the full range of reasons that a controller can rely on when using people's personal information to make significant automated decisions about them.¹ **REGULATOR**

The restriction was relaxed and not removed. Article 22B preserves the old position for special category data, where a solely automated significant decision remains prohibited unless the data subject gives explicit consent or the decision is necessary for contract performance and Article 9(2)(g) applies. Article 22C then sets four safeguards that apply in every case. The controller must provide information about the decision, enable representations to be made, enable human intervention to be obtained, and enable the decision to be contested.⁸ **STANDARD** A firm that reads section 80 as deregulation and stands down its human review and contest mechanisms has misread the Act.

4.1 The phrase that decides everything

Article 22A provides that a decision is based solely on automated processing if there is no meaningful human involvement in the taking of the decision.⁸ **STANDARD** That phrase is new to the statute. It does not appear in the old Article 22, and the Act supplies no definition, though Article 22A(2) directs that anyone considering the question must have regard, among other things, to the extent to which the decision is reached by means of profiling. Article 22D gives the Secretary of State power to define the term by affirmative-resolution regulations, and those regulations have not been made.⁸ **STANDARD**

So the test that determines whether the special category prohibition in Article 22B applies at all, and whether the Article 22C safeguards attach, currently has no authoritative meaning. The ICO consulted on draft guidance on automated decision-making and profiling between 31 March and 29 May 2026, and is under a statutory duty arising from secondary legislation under the Act to produce an AI code of practice. Neither the final guidance nor the code had been published as at the date of this paper.¹ **REGULATOR**

What that means operationally is that every firm is currently self-determining where meaningful human involvement begins. A claims handler who reviews a model's recommendation for four seconds and clicks accept may or may not constitute meaningful involvement. An underwriter who sees a referral queue of 200 cases a day and overturns three of them may or may not. My own view, which goes beyond anything the ICO has published, is that the even-

THE PRACTICAL POSITION

Controllers are setting the threshold themselves, in a period when the regulator has said what it thinks but has not yet said it in a form anyone can rely on.

tual definition will turn on whether the human had the information to judge, the authority to differ, and enough time to do either. Instrument all three now, and add a fourth measure for whether the involvement changed anything. Override rates and review durations are cheap to log and impossible to reconstruct.

4.2 What the courts have said about scoring

In December 2023 the Court of Justice of the European Union decided *OQ v Land Hessen*, generally known as *SCHUFA*. The Court held that the automated establishment by a credit information agency of a probability value concerning a person's ability to meet future payment commitments is itself a decision within Article 22(1), where a third party to whom that value is transmitted draws strongly on it in deciding whether to enter into a contract.³⁴ **CASE LAW** The argument that scoring is merely a preparatory step was rejected.

The practical effect in EU law is that Article 22 obligations attach to the scoring bureau and not only to the lender. As a CJEU judgment handed down after the transition period it does not bind UK courts, though it may be considered, and its relevance here is twofold. UK groups with EU operations must meet the supervisory practice it has shaped. And it poses, in EU terms, exactly the interpretive question that Article 22A now poses in UK terms, which is how much reliance on a machine output makes a decision effectively the machine's.

I would treat the reasoning as a reasonable guide to where the UK will land, while noting that the statutory wording is different and that a UK regulator has the freedom to reach a different answer. That is a prediction rather than a legal position, and a firm with real exposure on the point should take advice on it.

4.3 A control model that survives the definitional gap

Since the definition is pending, the sensible design does not depend on it. Build so that the answer to “was there meaningful human involvement” is yes on the evidence, whatever threshold is eventually set. Four things do most of that work, and none of them is expensive if designed in at the start.

EVIDENCE THAT A HUMAN WAS INVOLVED

01 Information The reviewer saw the material facts, the model's output, and the principal reasons for it, in a form they could act on. Logged as what was displayed, not what was available.	02 Authority The reviewer could reach a different conclusion without seeking approval. Recorded as the override right, and tested by whether overrides actually occur.	03 Time Review duration is measured and monitored. A median review time of a few seconds across a large queue is evidence against involvement.	04 Outcome Override rate, direction and downstream result are tracked by cohort, so the firm can show whether human involvement changed anything.
--	--	--	---

FIGURE 2 Four things to instrument before the definition arrives.

Proposed framework. Article 22A introduces the meaningful human involvement test without defining it, and Article 22D leaves the definition to regulations that have not been made. These four measures are my own suggestion for building evidence that will satisfy most plausible definitions. **ECAVEO**

The fourth is the one firms resist, because a low override rate is uncomfortable to look at. It is also the most informative. An override rate near zero tells you either that the model is excellent or that the human is rubber-stamping, and the two are distinguishable only by looking at the cases where the model was wrong and asking whether anybody caught it. That analysis needs a sample of known-bad cases seeded into the queue, which is a design decision nobody makes retrospectively.

A photograph of three people in a professional setting. On the left, a man with glasses and a beard is gesturing with his hands while speaking. In the center, a woman with curly hair is listening attentively with her hand on her chin. On the right, a man with a beard is also listening. They are seated around a dark wooden table with a laptop, papers, and coffee cups. In the background, there is a large window with a view of a city skyline, including a prominent church spire. A bookshelf with books and a lamp is also visible.

PART TWO · THE PORTFOLIO

Six workflow families, ranked by how confidently the value can be shown

CHAPTER FIVE

Where the value is credible

A portfolio of use cases is a set of hypotheses. Ranking them by demonstrable value produces a different order from ranking them by how interesting they are.

What follows is a portfolio, and the word is meant seriously. Each entry is a hypothesis about where machine assistance improves a workflow, and each carries a test that would falsify it. The ordering reflects how confidently the value can be demonstrated with the evidence that exists in September 2026, which is a different ordering from the one most strategy decks use.

Two general findings from Chapter 1 shape the whole portfolio. The measured gains concentrate among less experienced staff, which means the business case is stronger where a firm has high turnover or a wide skill distribution and weaker in a small team of specialists. And the gains reverse on tasks that sit outside the system's competence, which means every deployment needs a defined boundary and a way of knowing when work has crossed it.

THE SECTOR PORTFOLIO

WORKFLOW	OPPORTUNITY	VALUE TEST
Financial crime and fraud	Prioritise alerts, connect entities across accounts and products, and draft evidence-linked investigation narratives that a human closes.	Recall against known cases, precision, investigator time per case, and evidential traceability at the point of a suspicious activity report.
Claims	Triage documents and images, identify fraud signals, draft communications, and rank cases for human attention.	Cycle time, leakage, claimant outcomes by cohort, false positive rate, and results on appeal.
Insurance distribution and service	Summarise customer needs, explain policy language in plain terms, support vulnerable customers, and route complex cases to a person.	First-contact resolution, comprehension testing, complaint volume, accessibility, and the rate of harmful errors.
Wealth advice workflow	Prepare meetings, capture objectives, retrieve approved research, and assemble suitability evidence for an adviser to sign.	Adviser time per client, evidence completeness in the suitability file, client understanding, and corrections made at review.
Underwriting and pricing	Enrich risk assessment, detect missing evidence, and support underwriter review of referred cases.	Calibration, stability across renewal cycles, fairness by cohort, override quality, and portfolio outcomes over a full year.
Risk, finance and operations	Support scenario analysis, controls testing, regulatory reporting and retrieval of internal operational knowledge.	Timeliness, reconciliation exceptions, resilience under stress, and total cost including assurance.

TABLE 2 Six workflow families. The value test states what would have to improve against a real baseline, since a benefit nobody measured against a baseline will not survive its first budget review.

5.1 Why financial crime sits at the top

Financial crime investigation has the best structural fit of anything in this list, for reasons that have little to do with the technology. The work is document-heavy. The output is a narrative supported by evidence. There is a large corpus of historical cases with known outcomes, which means an evaluation set exists without anybody having to build one. And the cost of a false positive is measured in investigator hours, which the firm already counts.

The scale justifies the attention. UK Finance reported total payment fraud losses of £1.28bn in 2025, up 4 per cent, comprising £703.4m of unauthorised fraud across 3.81 million cases and £576.4m of authorised push payment fraud across 248,070 cases. It also reported £1.68bn of unauthorised fraud prevented, which it renders as 70 pence stopped per £1 attempted, and £354.3m returned to APP victims at a 61 per cent reimbursement rate.³⁵ **VENDOR RESEARCH** These are member-reported figures from the banking trade body, which makes them authoritative for UK payment fraud and not independent.

The same report describes criminals using AI to scale and refine attacks, craft impersonation scams and deploy deepfakes during account openings, with synthetic identities and selfie injections bypassing identity checks, and calls agentic AI the next force multiplier for global criminals.³⁵ **VENDOR RESEARCH** Treat that as qualitative assertion rather than measurement, because no attribution methodology is published. It is still the clearest statement from inside the industry that the attack side is moving.

What I would not do is let the arms-race framing drive the design. The reason to build here is the evaluation set, not the threat.

5.2 Why underwriting sits near the bottom

Underwriting is where the commercial appetite is strongest and where the demonstrable value is weakest, which is an awkward combination. The appetite is understandable, since pricing accuracy compounds across a portfolio and small improvements in loss ratio are worth a great deal.

Three things hold it back. The feedback loop is long, because a pricing decision made today is validated by claims experience over a year or more, so a pilot cannot report a value result inside a normal evaluation window. The fairness exposure is high, and section 3.2 sets out why removing protected variables does not address it. And for life and health insurance the EU AI Act treats pricing as an Annex III high-risk category, which for a firm whose output reaches the Union brings conformity assessment obligations from 2 December 2027.² **STANDARD**

None of that makes underwriting a prohibited use. It makes it a use where a firm should expect to spend more on assurance than on the model, and where the honest first deliverable is a cohort measurement capability, with any pricing improvement arriving later if it arrives. Firms that build the measurement first tend to discover things about their existing pricing that have nothing to do with AI, which is uncomfortable and useful.

5.3 The advice perimeter, which is about to move

Roughly 9 per cent of UK adults access regulated financial advice.¹⁹ **ADVISORY** That gap is the strategic prize in retail financial services and it is why the Mills Review matters more than its length suggests. The adoption plan of July 2026 makes 10 recommendations, one of which is an FCA review of the consumer, competition and wider effects of advice-like outputs from large language models, whose findings are to inform HM Treasury's consideration of the regulatory perimeter.¹⁹ **ADVISORY**

The Mills Review reported that around 26 per cent of consumers trust general-purpose tools such as ChatGPT, Claude or Gemini for financial advice.¹⁷ **REGULATOR** Whatever a firm decides to build, its customers are already asking somebody else, and that somebody is not authorised, not covered by the Consumer Duty and not within the ombudsman's reach.

I would treat the perimeter review as the single most consequential live policy item for retail firms, and I would not build a delegated-advice capability ahead of it. Building the evidence layer underneath one is a different matter, and that work is safe in any scenario.

CHAPTER SIX

The model you did not know you had

A third of use cases are somebody else's implementation, an unknown further number arrived as a feature in software the firm already licenses, and validation cannot reach what discovery has not found.

The Prudential Regulation Authority's supervisory statement SS1/23 is the closest thing UK financial services has to a model risk rulebook, and its scope is narrower than most commentary suggests. It applies to UK-incorporated banks, building societies and PRA-designated investment firms that hold internal model approval for regulatory capital under the internal ratings based, internal model or internal model method approaches.⁴ **REGULATOR** It does not apply to insurers. It does not apply to banks without internal model permissions. A wealth manager is nowhere near it.

Within scope, it reaches further than people expect. It covers all models used to inform business decisions, includes vendor models, and expressly captures artificial intelligence in modelling techniques such as machine learning.⁴ **REGULATOR** Its five principles are model identification and risk classification, governance, model development and implementation and use, independent model validation, and model risk mitigants.

The statement was updated on 23 April 2026, and the update is minor. Through the low impact amendments instrument LIAF01/26 the PRA clarified that SS1/23 expectations are not conditions for internal model approval and gave firms newly granted permission 12 months to comply. The PRA characterised this as a clarification and made it without consultation.⁴ **REGULATOR** Anybody presenting the April 2026 update as an AI-driven revision of model risk policy is overreading it.

6.1 Independent validation, and the objection from the firms

Principle 4 requires independent model validation. In its conventional form that means a second team, outside the first line, reproducing the model's logic and satisfying itself that inputs map to outputs in a defensible way.

The Bank of England held three roundtables with regulated banking and insurance firms in late 2025, publishing a summary on 16 February 2026. Participants endorsed the principles-based approach as giving sufficient space to innovate. They also said something more interesting, which is that validating how inputs mapped to outputs was not tenable or fully effective for increasingly complex AI models, and that the emphasis needed to shift towards outcome-focused monitoring and guardrails.³⁶ **REGULATOR**

That is a direct challenge to Principle 4 as it is usually practised, raised by the regulated population to its regulator. The PRA has said something adjacent, having published supervisory observations in November 2025 on firms' compliance with SS1/23 in the context of machine learning, which is the closest thing to published UK supervisory practice on AI in model risk.³⁶ **REGULATOR** It does not answer the validation question. It is not obviously wrong. Reproducing the logic of a large model is genuinely infeasible in the way it was feasible for a logistic regression. What replaces it, though, is unresolved, and outcome-focused monitoring is a phrase that covers a great deal of ground between rigorous statistical process control and a dashboard nobody reads.

AN OPEN QUESTION

If firms cannot validate the mapping, and the supervisor accepts that, what exactly is Principle 4 asking for in an AI context? Neither party has answered this in public.

The same roundtables raised something that boards should take more seriously than the validation point. Participants observed that as AI models become embedded in agentic systems throughout core business processes, substituting between AI providers may become more difficult.³⁶ **REGULATOR** Lock-in is deepening at the model layer while the regulatory response has been built at the infrastructure layer, which brings us to the next section.

6.2 Concentration, and the half of it that has been addressed

The critical third parties regime came into force on 1 January 2025 under powers in the Financial Services and Markets Act 2023, with HM Treasury designating and the regulators overseeing.²¹ **REGULATOR** On 10 July 2026 the first designations were announced, effective 13 July, covering Amazon Web Services EMEA, Google Cloud EMEA, Microsoft Ireland Operations and Oracle Corporation UK.

Designated providers must identify and manage risks to their critical services effectively and maintain open and timely communication with regulators and the firms that rely on them, particularly during major incidents. The Bank was careful to state that designation is not the same as authorisation, and that oversight is limited to the resilience of the services provided to UK financial firms.²¹ **REGULATOR**

All four designations are cloud and infrastructure providers. No AI model provider has been designated. The Treasury Committee's recommendation was to designate major AI and cloud providers, and only the cloud half has been delivered.^{18,21} **REGULATOR** Set that against the 2024 survey, which found the top three providers accounted for 73 per cent of named cloud providers, 44 per cent of models and 33 per cent of data.⁷ **REGULATOR** The 73 per cent is now inside a supervisory regime. The 44 per cent is not.

The Financial Policy Committee reached the same territory from a different direction in its July 2026 report, assessing AI risk across four channels covering use in core decision-making, use in markets, operational risk from AI service providers, and the changing cyber threat environment, and judging that risks were likely to increase. On concentration it noted that to the extent firms come to rely on a handful of frontier AI providers to improve their defensive cyber capabilities, this could represent a new channel of concentration risk.³⁷ **REGULATOR**

6.3 Finding what you already run

Before any of the above becomes actionable, a firm has to know what it has. Discovery is unglamorous and it is where I would spend the first month of any engagement in this sector.

Three categories go missing consistently. Embedded features are the largest, meaning AI capability switched on inside a licensed product the firm already uses, which nobody procured as AI and nobody recorded as a model. Vendor implementations are the second, at a third of use cases by the 2024 survey's count, where the firm owns the outcome and the supplier owns the mechanism. Departmental tools are the third, built by capable people in the business, often genuinely useful, and generally unknown to the second line.

The discovery method that works is not a questionnaire. It is a combination of the licence estate, the expenditure ledger and a walk round the floor, cross-checked against what people describe themselves as doing. Questionnaires find what people know they should declare.

THE DISCOVERY-TO-ASSURANCE SEQUENCE

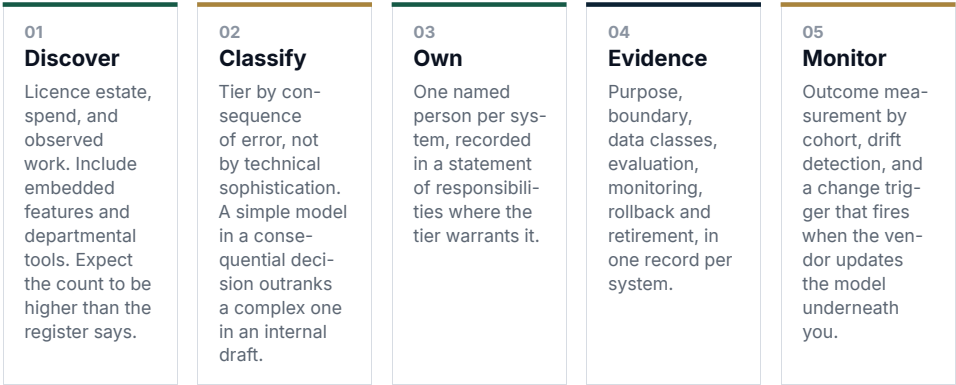


FIGURE 3 Five steps, in this order.

Proposed sequence. **ECAVEO** Classification by consequence rather than sophistication is drawn from the four-tier model used across this series, and the change trigger reflects the third-party dependency finding in the Bank of England and FCA survey of November 2024.⁷

Step five is the one that catches firms out. A vendor who improves their model has, from the firm’s point of view, deployed an unvalidated system into a live regulated process without notice. Contractual change notification is the control, and it has to be negotiated before signature because no supplier grants it afterwards.

CHAPTER SEVEN

Three deployments, and how to evaluate them

Each of these is bounded, reversible and evaluable inside a quarter. The evaluation protocol is written before anybody sees the tool, which is the part most programmes skip.

Three pilots, chosen because they cover different risk shapes and because each produces evidence that transfers to the next thing the firm builds. None of them makes a decision. All of them can be switched off on a Friday afternoon without anybody noticing on Monday.

7.1 The adviser evidence assistant

Scope. Permissioned meeting preparation and source-linked draft notes for a wealth adviser. The system retrieves from approved research and the client's own file, assembles a briefing, and drafts a file note after the meeting. It gives no advice and makes no suitability determination.

Why this one. The FCA's wealth data shows 13 per cent of firms using AI tools as at the end of 2024, so the sector is early and the competitive risk of moving is low.¹⁰ **REGULATOR** Suitability files are a known weak point in supervisory findings, the improvement is measurable, and a draft note reviewed by the adviser who attended the meeting has a natural verification step that costs nothing extra.

Evidence to collect. Completeness of the suitability evidence against a rubric agreed before the pilot. Count of unsupported claims in draft notes, meaning statements the system produced that are not traceable to an approved source. Adviser correction rate and correction type. Client comprehension, tested by asking a sample of clients to state in their own words what was recommended and why. Permission tests, meaning deliberate attempts to retrieve material the adviser is not entitled to see.

Stop condition. Any unsupported claim reaching a client file. Any permission leak at all.

7.2 The claims triage copilot

Scope. One bounded claim type, classified in shadow mode, with the system explaining the evidence behind each routing suggestion. Human handlers do not see the output during the shadow period.

Why this one. Claims volumes are large enough to give statistical power quickly. The Association of British Insurers reported £6.1bn paid across property claims in 2025, the highest annual total on record. Within that, more than 560,000 home insurance claims were paid at an average of £6,000, up 15 per cent year on year, so the wider total covers commercial property as well.³⁸ **VENDOR RESEARCH** Shadow mode means the pilot cannot harm a claimant, which removes the main objection to running it at scale.

Evidence to collect. Handling time against the pre-pilot baseline. Missed priority cases, which is the measure that matters and which only shadow mode can produce honestly. Fairness by cohort, defined and instrumented on day one. Override rate and direction once the system goes live. Appeal outcomes on cases the system touched, compared with matched cases it did not.

Stop condition. Any cohort disparity that the firm cannot explain by risk. Not cannot justify. Cannot explain.

7.3 The financial crime investigator

Scope. Summarise alerts and construct an evidence graph across accounts and products. The system does not close a case, does not escalate, and does not file anything.

Why this one. Section 5.1 sets out the structural fit. The historical case corpus gives a real evaluation set, and known-case recall is a measurable quantity in a way that most AI benefit claims are not.

Evidence to collect. Recall against a held-out set of known cases the system has not seen. False positive rate at the working threshold. Investigation time per case. Evidential traceability, meaning whether every assertion in a generated narrative resolves to a source record that an investigator can open.

Stop condition. Any narrative assertion that cannot be traced to a source.

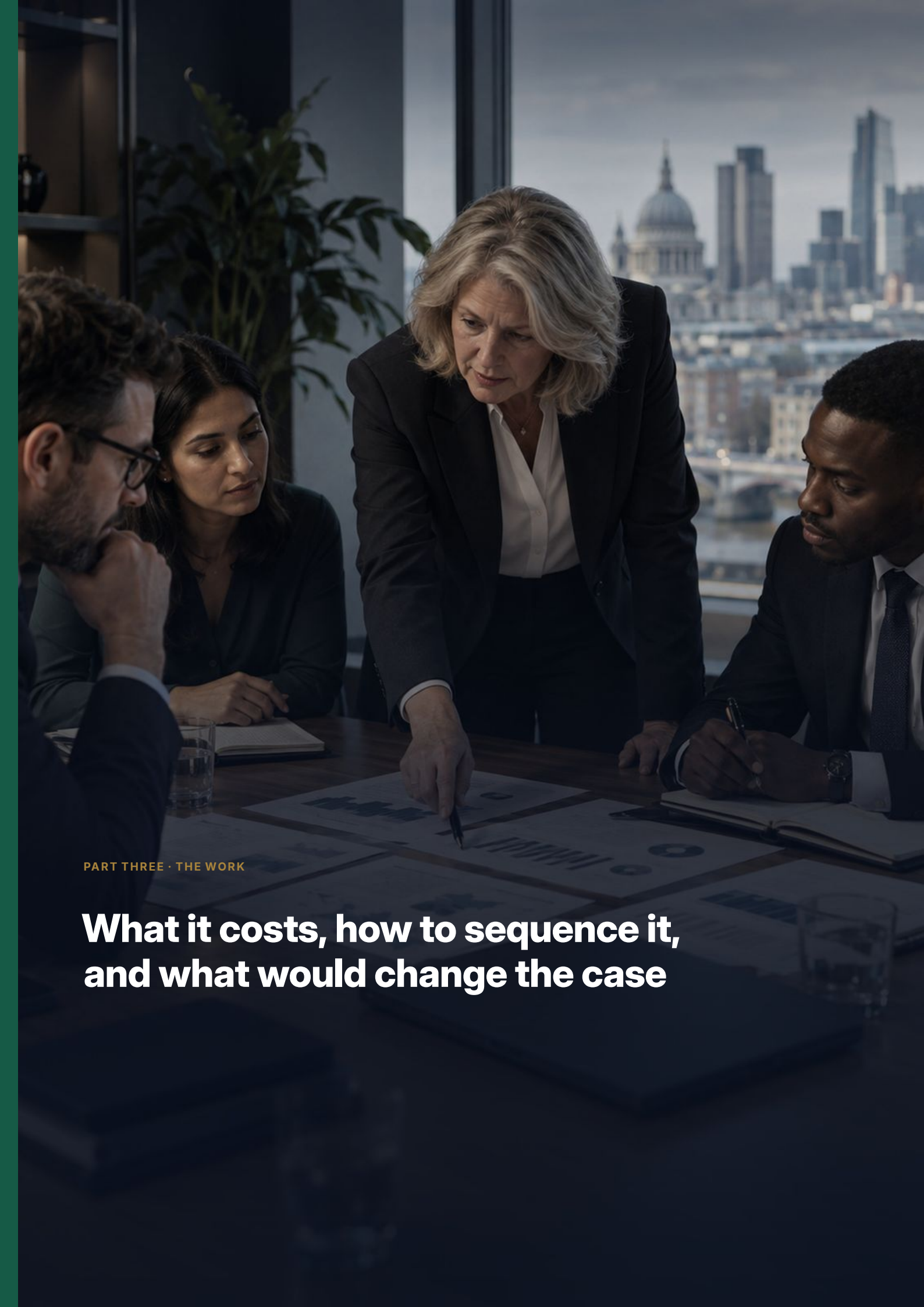
ILLUSTRATIVE SCENARIO

Northgate Mutual, a general insurer, runs the claims pilot

An invented mid-sized insurer runs the shadow-mode triage pilot on escape-of-water claims for 11 weeks. The model's routing agrees with the human handler in 84 per cent of cases. Of the 16 per cent where they differ, review finds the model right in about half.

The finding that changes the programme is elsewhere. Cohort analysis shows the model routing claims from one postcode group to enhanced validation at nearly twice the rate of the portfolio average. Risk-based explanations account for part of it. Part of it turns out to be a historical pattern in the training data reflecting an old and since-abandoned referral practice.

The pilot is stopped, the training set is rebuilt, and the second run does not show the disparity. Total cost of the finding is one quarter. Had the system gone live without cohort measurement, the finding would have arrived through the ombudsman, and the cost would have been the remediation of every claim it touched.



PART THREE • THE WORK

**What it costs, how to sequence it,
and what would change the case**

CHAPTER EIGHT

The economics of verification

The cost that decides whether a deployment pays is not the licence. It is the time somebody spends checking the output, and that cost is a design property.

Most business cases in this sector compare a licence fee against an estimate of time saved. Both numbers are usually wrong, and the error in the second one runs in a consistent direction.

Take the licence side first. The visible cost is per seat. The invisible costs are integration, identity and permissions work, evaluation, monitoring, the second-line review the tier attracts, incident handling, and the internal time spent answering questions from three different assurance functions. In the deployments I have worked on, the licence is rarely more than a quarter of the first-year cost and sometimes less than a tenth. That is a practitioner observation from a handful of engagements. Firms should build their own. **PRACTICE**

The time-saved side is worse, because of the perception problem from Chapter 1. When 16 experienced developers were slowed by 19 per cent and believed they had been sped up by 20 per cent, the error was not small and it did not correct itself on reflection.³⁰ **PEER-REVIEWED** Self-reported time saving is the standard measure in this sector's pilots. It is also, on the only evidence available, an unreliable one.

8.1 Verification cost as a design variable

Here is the part that changes decisions. The cost of checking an output is not fixed by the task. It is set by how the system presents its work, and a firm buying tools has a great deal of influence over it.

An output with a source link attached converts verification from re-performance into spot-checking. A reviewer opens the cited record, confirms the assertion, and moves on. Without the link, the reviewer has to reconstruct the answer independently to know whether it is right, which costs as much as doing the work and sometimes more, because reading somebody else's reasoning before forming your own is slower than starting cold.

WHAT VERIFICATION COSTS, BY DESIGN CHOICE

DESIGN FEATURE	WITH SOURCE LINKS AND ABSTENTION	WITHOUT
How a reviewer checks	Opens the cited record and confirms the assertion	Reconstructs the answer to know whether it is right
Marginal cost of a check	Seconds per assertion	Comparable to doing the task
What happens under time pressure	Spot-check a sample, with the rest traceable if challenged	Checking is dropped and the output is accepted
Evidence available six months later	The source trail persists in the record	Nothing beyond the output itself
Effect on the business case	Benefit survives the review burden	Benefit is consumed by it, or the review stops

TABLE 3 Illustrative comparison of the same task under two designs. The figures are structural rather than measured, and the point is the ratio between the columns. **ECAVEO**

Abstention is the second design property and it is rarer. A system that declines to answer when its retrieval returns nothing relevant is worth substantially more than one that answers anyway, because the failure mode changes from a plausible wrong answer to a visible gap. Ask a supplier whether their product abstains, and what its abstention rate is on your own material. Most cannot tell you, and the answer usually arrives as an explanation of why the question is difficult.

8.2 What the sector can afford, and what that implies

UK financial services contributed £223.6bn to the economy in 2025, 8 per cent of total output, employing 1.09 million people in financial and insurance services as at the fourth quarter of 2025.³⁹ **REGULATOR** The investment management industry reported £11.1 trillion of assets under management as at the end of December 2025, up 11 per cent, from 57 questionnaire responses representing 80 per cent of total assets, which makes it a trade body survey of its own members and reliable for aggregates while not independent.⁴⁰ **VENDOR RESEARCH**

Numbers of that size make almost any AI programme look cheap, which is precisely the problem. A cost that is immaterial at group level is still capable of producing no return, and the absence of any independent measurement of realised return in this sector means nobody can currently distinguish the programmes that are working from the ones that are merely affordable. **EVIDENCE GAP**

The discipline I would impose is a baseline, taken before anything is deployed, on a metric the business already reports. Not a new metric invented for the pilot. If the workflow has no existing measure, that is itself informative and probably means the workflow is the wrong place to start.

8.3 Controls that are not optional, whatever the tier

Six technical controls apply to every deployment in a regulated firm, and they are consistent with ICO accountability expectations, the National Cyber Security Centre's secure AI development guidelines of November 2023, and the NIST generative AI risk profile of July 2024.^{1,3,13} **STANDARD** The NIST profile is a United States voluntary framework with no UK legal force, and it transfers only as a control taxonomy that firms adopt by choice. Firms wanting a certifiable management system around the same ground have ISO/IEC 42001, published in December 2023.¹⁴ **STANDARD**

- Enterprise identity and client, matter or engagement permissions are enforced before retrieval and before any tool call, and not applied to results afterwards.
- Contracts define training use, retention, hosting region, subprocessors, deletion, security incident notification, model change notification and audit rights.
- Agents operate with least-privilege, short-lived credentials, and any consequential action requires explicit approval and is reversible.
- Every production system records the model, prompt, retrieval corpus, tools, permissions and evaluation version behind each output.
- Prompt injection, malicious documents, sensitive-data disclosure and insecure output handling are tested before scale, not after.
- Generated work passes the same review and record-keeping controls as work produced without AI.

The first is the one that gets skipped, because filtering after retrieval is easier to build and looks identical in a demonstration. It is not identical in an incident, and the difference shows up in the logs.

CHAPTER NINE

The first 90 days, then the rest of the year

A sequence that produces evidence before it produces exposure, and that leaves a firm better placed whatever the regulator does next.

9.1 The first 30 days

Interview leaders and observe real work across the priority functions, which for most firms means claims, advice and financial crime. Watching is more informative than asking, and the gap between the two is where the use cases are.

Inventory approved and shadow AI use, including capability embedded in software the firm already licenses. Section 6.3 sets out why the questionnaire method fails.

Map data boundaries, contractual duties, permissions and the individuals who currently own each regulated decision. This is the input to the accountability allocation and it usually reveals that two people each believe the other owns something.

Name the accountable senior manager and amend their statement of responsibilities. Since no prescribed responsibility exists, the firm is making its own allocation, which is both the difficulty and the opportunity.^{5,18} **REGULATOR**

Publish interim safe-use guidance and a simple route for proposing an experiment. Firms that skip this get shadow adoption and then have to run the discovery exercise twice.

Select three pilots with baselines, representative test sets and stop conditions written down before anyone sees a tool.

9.2 Weeks 5 to 12, and then months 4 to 12

THE 12-MONTH SEQUENCE

PHASE	TIMING	OUTPUTS
Frame	Weeks 1 to 4	Consumer journeys, regulated decisions, model and third-party inventory, outcome baselines, named accountability.
Prove	Weeks 5 to 12	Three shadow or assisted pilots with cohort testing, held-out evaluation sets and independent challenge from outside the delivery team.
Assure	Months 4 to 6	Inventory integrated with model validation where SS1/23 applies, Consumer Duty outcome evidence, privacy, security and resilience assessments, contractual change triggers.
Scale	Months 7 to 12	Proven patterns extended by product and channel, with monitoring for outcomes, drift, concentration and appeal results.

TABLE 4 Each phase produces evidence the next phase depends on. Phases can overlap, and the ordering of the evidence cannot.

Independent challenge in the Prove phase is the item most often cut and the one I would defend hardest. The person who built the pilot cannot evaluate it, not because of dishonesty but because they know which cases it handles well and will unconsciously sample them. A colleague from a different function, given the test set and half a day, finds things the builder cannot see.

9.3 What success looks like after nine months

Scaled workflows with measured quality, time or service improvement against the baselines set in month one. A single intake route, approved systems, documented use cases and monitored supplier changes. Users who can explain what a system cannot do and escalate a failure when they meet one. Business owners who hold the controls and the benefits after the programme team has moved on. A backlog that separates proven opportunities from research and from uses the firm has decided against.

That last item is the one that signals maturity. A portfolio with nothing in the decided-against column has not been governed, it has been approved.

CHAPTER TEN

Where the evidence is thin

Everything in this paper that would change if the underlying evidence changed, stated plainly, so a reader can judge how much weight the argument will bear.

10.1 What I could not establish

No randomised or quasi-experimental study measures AI effects in insurance underwriting, claims handling or regulated financial advice. The productivity evidence in Chapter 1 comes from customer support, consulting and software development, and its application here is analogy. **EVIDENCE GAP**

No independent official statistic for AI adoption in UK financial services exists, because the survey that produces the UK's official adoption figures excludes finance and insurance from its coverage.²² **REGULATOR** Every figure in circulation is a regulator surveying its own population, a trade body counting members, or vendor research. **EVIDENCE GAP**

I could not establish whether the Information Commissioner has taken enforcement action against a financial services firm arising from AI, algorithmic pricing or automated decision-making. The enforcement register loads its records through a mechanism I could not query reliably, so the honest statement is that I did not find any, which is different from there being none. **EVIDENCE GAP**

No AI model provider is known to be under consideration for critical third party designation. The absence of information is not evidence that none is. **EVIDENCE GAP**

The Financial Stability Board's sound practices for responsible adoption of AI, consulted on between 10 June and 22 July 2026, had not been published in final form.²⁶ **REGULATOR** The 12 proposed sound practices span the AI lifecycle, covering generative and agentic systems. They sit in a consultation and should not be described as standards.

10.2 What would change the argument

Publication of FCA guidance on senior manager accountability for AI. This is the single item that would most change the paper. If a prescribed responsibility or an equivalent is created, the central argument in Chapter 2 becomes a historical note and the sequencing in Chapter 9 becomes compliance rather than anticipation.

A 2026 or 2027 edition of the joint Bank of England and FCA survey. If comprehension has improved materially against the 46 per cent partial understanding figure, the title of this paper stops being accurate.⁷ **REGULATOR**

The Secretary of State's regulations defining meaningful human involvement under Article 22D, or final ICO guidance. Chapter 4 is written around an undefined statutory term, and a definition would replace judgement with a test.⁸ **STANDARD**

Designation of a model provider as a critical third party. That would close the gap identified in Chapter 6 and would substantially change how firms should think about substitution risk.

Any rigorous evaluation of AI in underwriting or claims. One well-designed study would do more for this sector's decision-making than another year of surveys.

10.3 Where I have gone beyond the sources

Three claims in this paper are mine. I would defend each of them, while accepting that a reasonable reader might land somewhere else.

The first is that the absence of a prescribed responsibility for AI is the load-bearing weakness in the FCA's approach, and that the argument about who should have owned a failure will be resolved against whichever individual held the operational function. Neither the FCA nor the Treasury Committee says this. **PRACTICE**

The second is that the eventual definition of meaningful human involvement will turn on information, authority, time and demonstrated effect. That is a prediction about a definition nobody has drafted. **PRACTICE**

The third is the verification economics in Chapter 8. The claim that source links and abstention determine whether a deployment pays is a structural argument supported by practitioner experience and not by any published measurement I could locate. **ECAVEO**

There is a fourth thing I have not resolved and will not pretend to have. The EIOPA opinion of 6 August 2025 sets supervisory expectations for insurance AI governance, and it deliberately excludes systems that are prohibited or high-risk under the EU AI Act, addressing those of limited and minimal risk so as to avoid regulatory overlap.¹¹ **STANDARD** For a UK insurer with EU operations that produces an odd division, where the everyday tail of AI is covered by one set of expectations and the consequential pricing models by another. Whether that division is workable in practice is a question I have heard raised in industry discussion and have not seen answered anywhere, and I do not have a view worth committing to print.

References

SOURCES, WITH METHOD AND PROVENANCE WHERE IT BEARS ON THE WEIGHT A CLAIM CAN CARRY

Important claims were checked against regulators, professional bodies, standards organisations and primary official material. Where a source is a survey from a body with a commercial or membership interest in the result, the reference says so. Guidance and timetables in this area moved repeatedly during 2026 and remain subject to change.

1. Information Commissioner's Office. *Guidance on AI and data protection*. Last updated 15 March 2023, and carrying a notice that it is under review following the Data (Use and Access) Act. <https://ico.org.uk/for-organisations/uk-gdpr-guidance-and-resources/artificial-intelligence/guidance-on-ai-and-data-protection/>
2. European Union. *Regulation (EU) 2024/1689 (Artificial Intelligence Act)*, Article 2 and Annex III points 5(b) and 5(c), as amended by *Regulation (EU) 2026/1744 (Digital Omnibus on AI)* of 8 July 2026, in force 27 July 2026, deferring stand-alone Annex III obligations to 2 December 2027 and embedded high-risk obligations to 2 August 2028. https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ:L_202601744
3. National Cyber Security Centre. *Guidelines for secure AI system development*. Version 1.0, 27 November 2023, co-sealed with CISA and 21 other international agencies. <https://www.ncsc.gov.uk/collection/guidelines-secure-ai-system-development>
4. Prudential Regulation Authority. *SSI/23: Model risk management principles for banks*. Published 17 May 2023, updated 23 April 2026 through LIAF01/26. Applies to UK-incorporated banks, building societies and PRA-designated investment firms holding internal model approval; does not apply to insurers. <https://www.bankofengland.co.uk/prudential-regulation/publication/2023/may/model-risk-management-principles-for-banks-ssi>
5. Financial Conduct Authority. *AI and the FCA: our approach*. First published 8 September 2025, last updated 13 February 2026. <https://www.fca.org.uk/firms/innovation/ai-approach>
6. Smith, A., Head of Cross-cutting Policy and Strategy, Financial Conduct Authority. *AI in financial services: shaping our approach through industry engagement*. 8 June 2026. <https://www.fca.org.uk/news/blogs/ai-financial-services-approach>
7. Bank of England and Financial Conduct Authority. *Artificial intelligence in UK financial services 2024*. 21 November 2024. Joint structured questionnaire, 118 responding firms across six sectors; response rate not published; firm counts not weighted by size. <https://www.bankofengland.co.uk/report/2024/artificial-intelligence-in-uk-financial-services-2024>
8. *Data (Use and Access) Act 2025*, section 80, substituting Articles 22A to 22D for Article 22 UK GDPR. In force 5 February 2026 by SI 2026/82. <https://www.legislation.gov.uk/ukpga/2025/18/section/80>
9. Financial Conduct Authority. *CP26/23: Consumer Duty, scope and proportionality*. Published 29 June 2026, closing 18 September 2026. Contains no AI-specific provisions. <https://www.fca.org.uk/publications/consultation-papers/cp26-23-consumer-duty-scope-and-proportionality>
10. Financial Conduct Authority. *Wealth management survey report 2026*. Published 18 August 2026. Approximately 400 firms, 234 completing all three waves; data as at 31 December 2024. <https://www.fca.org.uk/data/wealth-management-survey-report-2026>
11. European Insurance and Occupational Pensions Authority. *Opinion on Artificial Intelligence governance and risk management*. 6 August 2025. Addressed to national competent authorities and not legally binding; expressly excludes systems prohibited or high-risk under the EU AI Act. https://www.eiopa.europa.eu/publications/opinion-artificial-intelligence-governance-and-risk-management_en
12. International Organization of Securities Commissions. *Supervisory Toolkit for AI Use in Capital Markets*, FR/02/2026, May 2026, superseding in part the consultation report CR/01/2025 of March 2025. <https://www.iosco.org/library/pubdocs/pdf/IOSCPD823.pdf>
13. National Institute of Standards and Technology. *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1. 26 July 2024. A United States voluntary framework with no UK legal force. No later revision exists. <https://doi.org/10.6028/NIST.AI.600-1>
14. International Organization for Standardization and International Electrotechnical Commission. *ISO/IEC 42001:2023, Artificial intelligence management system*. December 2023.
15. Financial Conduct Authority. *PS21/5: General insurance pricing practices market study, feedback to CP20/19 and final rules*. May 2021. Pricing remedy in force 1 January 2022. <https://www.fca.org.uk/publication/policy/ps21-5.pdf>

16. Financial Conduct Authority. *EP25/2: An evaluation of our General Insurance Pricing Practices remedies*. July 2025. Continuous difference-in-differences; 16 home insurers at c.80 per cent market share and 13 motor insurers at c.57 per cent; Q1 2019 to Q1 2024. <https://www.fca.org.uk/publication/corporate/ep25-2.pdf>
17. Financial Conduct Authority. *AI and the future of retail financial services (the Mills Review)*. Led by Sheldon Mills, published 6 July 2026. Consumer research by Yonder Consulting, April 2026, 5,000+ UK retail financial services consumers, quota-controlled; commissioned and funded by the FCA. <https://www.fca.org.uk/publications/corporate-documents/mills-review>
18. Treasury Committee. *Artificial intelligence in financial services*. Fifteenth Report of Session 2024 to 2026, HC 684, 20 January 2026. <https://publications.parliament.uk/pa/cm5901/cmselect/cmtreasy/684/report.html>
19. Rees, H. and Dhawan, R. *Financial Services AI Adoption Plan*. Published by HM Treasury, 14 July 2026. An independent report by the government's AI Champions for financial services, not a statement of HM Treasury policy. Ten recommendations, the first of which is that regulators consider confirming the principles-based, outcomes-focused approach. <https://www.gov.uk/government/publications/ai-adoption-plan-financial-services/financial-services-ai-adoption-plan>
20. Treasury Committee. *AI in financial services: responses to the Committee's Fifteenth report*. Seventh Special Report of Session 2024 to 2026, 16 April 2026. <https://publications.parliament.uk/pa/cm5901/cmselect/cmtr easy/1791/report.html>
21. Financial Conduct Authority. *PS24/16: Operational resilience, critical third parties to the UK financial sector*, 12 November 2024, issued by the PRA as PS16/24, rules in force 1 January 2025; and Bank of England, first designations announced 10 July 2026, effective 13 July 2026, covering Amazon Web Services EMEA, Google Cloud EMEA, Microsoft Ireland Operations and Oracle Corporation UK. <https://www.bankofengland.co.uk/news/2026/july/uk-financial-regulators-to-begin-overseeing-critical-third-parties-announced-by-hmt>
22. Office for National Statistics. *Artificial intelligence in UK businesses: 2023 to 2026*. 20 July 2026. Business Insights and Conditions Survey wave 159, fieldwork 15 to 28 June 2026, 38,637 responses, 26.7 per cent response rate. Finance and insurance are outside the survey's coverage. <https://www.ons.gov.uk/businessandtrade/business/businessservices/articles/artificialintelligenceinukbusinesses/2023to2026>
23. Financial Conduct Authority. *FS25/5: AI Live Testing*, feedback statement 9 September 2025 drawing on 67 responses; second cohort of eight firms announced 21 April 2026. Evaluation report due end 2026 or early 2027. <https://www.fca.org.uk/publications/feedback-statements/fs25-5-ai-live-testing>
24. Financial Ombudsman Service. *Annual complaints data and insight 2025/26*. 21 May 2026. The year-on-year fall follows the introduction of referral charges for professional representatives and should not be read as improved conduct. <https://www.financial-ombudsman.org.uk/businesses/resolving-complaint/our-insight/annual-complaints-data-and-insight-2025-26>
25. Harris, M., Chief Operating Officer, Financial Ombudsman Service. *Embracing AI's transformational impact on consumer complaints*. 27 March 2026; and the Service's response to the Mills Review, February 2026. The 35 per cent figure is described by the Service as early small-sample analysis. <https://www.financial-ombudsman.org.uk/businesses/resolving-complaint/our-insight/embracing-ais-transformational-impact-consumer-complaints>
26. Financial Stability Board. *Sound Practices for Responsible Adoption of Artificial Intelligence: consultation report*. 10 June 2026, closed 22 July 2026. A consultation, not a standard. <https://www.fsb.org/2026/06/sound-practices-for-responsible-adoption-of-artificial-intelligence-ai-consultation-report/>
27. Brynjolfsson, E., Li, D. and Raymond, L. *Generative AI at Work*. Quarterly Journal of Economics, volume 140, issue 2 (2025), pp. 889 to 942. 5,179 customer support agents; staggered rollout used for identification; single firm and single task type.
28. Cui, Z., Demirer, M., Jaffe, S., Musolf, L., Peng, S. and Salz, T. *The Effects of Generative AI on High-Skilled Work: Evidence from Three Field Experiments with Software Developers*. Management Science. Three randomised experiments, 4,867 developers, at Microsoft, Accenture and a third large firm. The tool studied is a Microsoft product and two of the three sites are the vendor and a systems integrator.
29. Dell'Acqua, F., McFowland, E., Mollick, E., and others. *Navigating the Jagged Technological Frontier*. Organization Science. Randomised field experiment, 758 consultants. BCG collaborated on the design and supplied the participants and tasks.
30. Becker, J., Rush, N., Barnes, E. and Rein, D. *Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity*. arXiv:2507.09089, 12 July 2025, revised 25 July 2025. Randomised, 16 developers, 246 tasks. Very small sample of self-selected experts on mature codebases; does not generalise. <https://arxiv.org/abs/2507.09089>
31. Citizens Advice. *Citizens Advice sounds the alarm on £280 car insurance ethnicity penalty*. 22 March 2022. Charity-commissioned observational research, 18,000 car insurance costs from debt clients plus 649 mystery-shopping exercises by Europe Economics. Cannot separate ethnicity from correlated geographic

- risk factors. <https://www.citizensadvice.org.uk/about-us/media-centre/press-releases/citizens-advice-sounds-the-alarm-on-280-car-insurance-ethnicity-penalty/>
32. Bartlett, R., Morse, A., Stanton, R. and Wallace, N. *Consumer-Lending Discrimination in the FinTech Era*. Journal of Financial Economics (2021); NBER working paper 25943. United States mortgage market, GSE-conforming loans; pre-LLM era. <https://www.nber.org/papers/w25943>
 33. Fuster, A., Goldsmith-Pinkham, P., Ramadorai, T. and Walther, A. *Predictably Unequal? The Effects of Machine Learning on Credit Markets*. The Journal of Finance (2022), DOI 10.1111/jofi.13090. A counterfactual simulation study using United States mortgage data, not an experiment.
 34. Court of Justice of the European Union. *OQ v Land Hessen*, Case C-634/21, First Chamber, judgment 7 December 2023. Not binding on UK courts post-transition, though it may be considered. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A62021CJ0634>
 35. UK Finance. *Annual Fraud Report 2026*. 15 June 2026. Member-reported data from the banking trade body. The attribution of attack scaling to AI is qualitative assertion rather than measurement. <https://www.ukfinance.org.uk/policy-and-guidance/reports/annual-fraud-report-2026>
 36. Bank of England. *Summary of AI roundtables*. Three roundtables with regulated banking and insurance firms held in late 2025, summary published 16 February 2026; and Prudential Regulation Authority, supervisory observations on firms' compliance with SS1/23 in the context of artificial intelligence and machine learning, published November 2025 following the model risk management roundtable of October 2025. <https://www.bankofengland.co.uk/minutes/2026/february/summary-of-ai-roundtables-feb-2026>
 37. Bank of England. *Financial Stability Report, July 2026*. 7 July 2026. <https://www.bankofengland.co.uk/financial-stability-report/2026/july-2026>
 38. Association of British Insurers. *Adverse weather pushes property insurance payouts to £6.1 billion in 2025*. 19 February 2026. Trade body data. The £6.1bn covers property claims overall; the 560,000 claims and the £6,000 average are the home insurance figures within it.
 39. House of Commons Library. *Financial services in the UK*, SN06193. Last updated 11 September 2026. <https://commonslibrary.parliament.uk/research-briefings/sn06193/>
 40. The Investment Association. *UK investment management industry hits record £11.1 trillion AUM*. 6 August 2026. Trade body survey of its own members; 57 questionnaire responses representing 80 per cent of total assets under management. <https://www.theia.org/news/press-releases/uk-investment-management-industry-hits-record-ps111-trillion-aum>

ECAVEO RESPONSIBLE AI SERIES

Partial Understanding

Ecaveo Responsible AI Series, Financial Services. Version 1.0, September 2026.

ALSO IN THIS SERIES

The Traceable Thesis. Private equity. Sufficient and Appropriate. Accounting and audit.

Verified Authority. Law firms. The Protected Constraint. Operations.

THE AUTHOR

Paul Forrest is a fractional Head of AI working with private equity and venture capital portfolio businesses on the deployment of artificial intelligence.