

Chartered
Insurance
Institute

Standards. Professionalism. Trust.

Responsible AI: from policy to practice

Roundtable summary report

Contents

As a chartered body with a public interest mandate, the Chartered Insurance Institute provides a forum where stakeholders can collaborate on shared challenges. Our independence enables honest dialogue, facilitating the development of sector-wide guidance and recommendations that strengthen professional standards and deliver better customer outcomes.

Foreword.....	3
Executive summary.....	4
Participants.....	5
Introduction.....	6
The starting point: purpose before technology.....	6
• Use case, measurement and objectives.....	6
• Enhancing processes versus redesigning them.....	6
What do we mean by “AI”? Definitions, the spectrum and trust.....	7
• Trust, friction and the invisible networks.....	7
Human judgement and the limits of “human in the loop”.....	8
• Cognitive offloading and accountability.....	8
• Augmented, not artificial, intelligence.....	8
The ethical challenge: from population fairness to the individual.....	9
Governance, procurement and operational resilience.....	10
Skills, capability and CPD.....	11
What the CII and other bodies can do.....	12
• Make corporate Chartered status a genuine value proposition.....	12
• Develop AI fluency and behavioural CPD.....	12
Conclusion and next steps.....	13
Appendix 1 Artificial Intelligence in Insurance and Personal finance – a policy scaffold.....	14
Appendix 2.....	16
Appendix 3 Further resources.....	18

Foreword

Chartered firms publicly commit to leading the profession in technical competence and ethical standards that deliver better customer outcomes. Those standards benefit customers, the profession and our standing in society. That matters even more as artificial intelligence (AI) becomes embedded in product design, decision-making and client support. AI creates real opportunities, but it also sharpens questions of judgement, accountability, fairness and trust.

This paper puts that debate where it belongs: at the heart of professional leadership. For firms committed to Chartered status, the question is not only whether AI can be used efficiently, lawfully or competitively, but also, depending on the use case and its potential impact on customers, whether its use reflects the spirit of the [CII Code of Ethics](#): acting with integrity, serving each client's best interests, delivering a high standard of service and treating people fairly. These considerations are central to being a trusted business.

Chartered status is a public declaration of a firm's commitment to public trust, investment in people and the development of the profession. As expectations of corporate Chartered evolve, responsible AI governance is becoming part of how firms show those commitments in practice, reinforcing customer confidence, ethical culture and sound professional judgement across the organisation.

This report summarises a roundtable discussion held in that spirit, focusing on ethical considerations, practical challenges and the areas where professional standards can shape better decisions.

Executive summary

On 4 June 2026, the Chartered Insurance Institute (CII) convened senior representatives from Chartered firms, technology, academia, trade and professional bodies to discuss what responsible AI adoption should look like in insurance and financial planning, and what role the CII should play. The session focused on three questions: what good AI governance looks like; how the CII can codify good practice through guidance, training and standards; and how it can support confident, responsible adoption that benefits customers and society.

One key message prevailed: start with purpose and the business problem, not the technology. Much current adoption is driven by fear of being left behind or by the need to show efficiencies rather than by defined goals tied to organisational purpose. Exploring AI because it may cut costs can be a reason to test it, but it is not a strategy. If the aim is growth, efficiency, better customer outcomes, or a mix of these, firms need to choose that aim explicitly and align governance to it.

Several tensions ran through the discussion. AI should be understood as a spectrum of techniques, not a synonym for large language models, because conflating them distorts both risk and opportunity. The assumption that AI automatically raises productivity was challenged, with examples ranging from complaint systems overwhelmed by AI-generated volume to firms quietly reversing automation. “Human in the loop” was seen as necessary but not enough: oversight must be active rather than a rubber stamp, and the better a system appears to perform, the greater the risk of cognitive offloading. The professional remains personally accountable for the outcome. Efficiency should strengthen human skills such as curiosity, empathy and critical thinking, not reduce people to quality control for automated processes.

On ethics, participants noted that most published AI policies frame fairness at population level, while professional ethics requires acting in the best interests of each individual customer. That shifts attention from model-level fairness to individual judgement, suitability and explanation, and brings conflicts of interest and the duty to challenge into sharper view. A recurring conclusion was that existing regulation and the CII Code of Ethics already cover much of this; the priority is applying them well, not drafting new rules.

Underpinning many of these themes was a more fundamental constraint: a critical AI fluency gap within the sector. Where understanding of what AI is, what it can and cannot do, and where its risks lie is limited, firms struggle to formulate robust policy, exercise effective governance and make sound adoption decisions. Closing that gap – in boardrooms, risk functions and front-line roles alike – is a precondition for everything else this report discusses.

Participants saw clear roles for the CII and other bodies: define what good looks like through a flexible framework rather than boilerplate; make corporate Chartered status a stronger value proposition, including a consumer-facing signal of responsible practice; develop AI fluency and behavioural CPD; provide maturity benchmarking and kitemarks or certification; and work with professional and trade bodies on shared governance and procurement standards.

Participants

The participants below contributed to the discussion, but this report does not represent the views of any individual or their employer.

Rebecca Aston	CII
Lauren Branston	Institute of Business Ethics
Kanika Chaganty	Brit
Matthew Connell	CII
Chris Cowton	University of Huddersfield Emeritus Professor and CII Professional Standards Committee
Scott Daniels	Plus Group
Neil Freshwater	CII Professional Standards Committee
Annabel Gillard	Conversationsonai.org
Adam Harper	CII
Sandeep Karkhanis	Allianz
Connor McNicholas	Magnetic
David Otudeko	ABI
Nick Onslow	The Progeny Group
Pam Quinn	BIBA
Ian Simons	CII

Introduction

AI is already playing a central and increasing role in serving customers and supporting professionals across insurance and financial planning, bringing both opportunities and risks in equal measure. As a Chartered body whose core purpose is to build and maintain public trust, the CII has both an interest and a responsibility in helping firms and professionals navigate this shift well.

The roundtable was convened to gather expert insight on four questions: what good AI governance looks like in this sector; where legal, compliance and ethical approaches leave gaps or create tensions for professional judgement; what role the CII should play through guidance, training and standards, including expectations of Chartered firms; and how it can support ethical, effective adoption through thought leadership that goes beyond warning about risk.

The starting point: purpose before technology

Many firms are adopting AI because they feel they need to be seen to act, bolting tools onto existing processes instead of asking the harder question first: where, if anywhere, does AI genuinely serve the organisation's strategic objectives?

Use case, measurement and objectives

A practical framework emerged for any proposed use case: define the use case, what is being measured and the objective. Participants distinguished between efficiency goals such as cost and speed, and customer-outcome goals, noting that speed may serve the firm, the customer or both. The discipline is to be explicit about the intended benefit and, if objectives clash, which takes priority. A third lens also matters: the business is social as well as financial, so firms should consider AI's effect on the workforce and on internal networks of trust, involving employees early rather than presenting change as a fait accompli.

Enhancing processes versus redesigning them

Participants also warned against bolting technology onto a failing process. Several argued for moving from risk-led to purpose-led policy: start with why the organisation exists, then ask which parts of that purpose AI is meant to serve, and what happens if it fails to deliver. As confidence grows, firms may move beyond redesigning existing processes to creating propositions that do not yet exist. That is why they will need principles, not just rules.

"The foundational action of applying AI responsibly is to work out whether you should be doing it to begin with – what it is, what we use it for, and what the risks are."

What do we mean by “AI”? Definitions, the spectrum and trust

Participants warned that the term “AI” has become so loose that it often obscures more than it clarifies. In practice it is increasingly used to mean large language models, or whichever label is in fashion, most recently “agents”. Yet AI covers a wide range of techniques, from optical character recognition and classical machine learning to graph neural networks used to model wildfire spread. What matters is to match the technique to the problem instead of throwing an LLM at it and hoping for the best. Shared definitions, such as the OECD’s, would help firms classify what they are using and reason more clearly about risk.

Trust, friction and the invisible networks

Participants observed that the relationships of trust and belief that underpin the sector have always been largely invisible. Trust was likened to an ecosystem: with visible markers, like trees in a forest, and complex networks of interdependent actors like mycelium, roots and bacteria that may not be visible or directly coordinated but coexist symbiotically to support a healthy environment. Disrupting that invisible ecosystem – even with good intentions – can lead to unintended systemic consequences, not least because no single actor within the ecosystem can control the consequence on all other dependants.

One participant said we may be the first generation that cannot be sure what we see or hear is real. “AI slop” — content generated, processed and reviewed by models with little human input — was seen as a corrosive risk. Underwriters, for example, may receive polished material from intermediaries without confidence that anyone has actually read or understood it.

That led to a challenge to a common assumption: less friction is not always better. Friction can protect judgement and support trust. A faster process may speed up information exchange while reducing insight, engagement and confidence. In some national-security settings, friction is deliberately built in so people cannot simply press go. Participants also worried about the flattening effect of shared platforms, drawing a parallel with music-streaming algorithms that homogenise output and weaken brand, culture and differentiation.

Human judgement and the limits of “human in the loop”

Participants strongly supported human oversight but were sceptical of “human in the loop” as a slogan. A person being involved does not guarantee a good outcome. What matters is active oversight: someone using judgement, analysis and critical thinking rather than passively signing off.

Cognitive offloading and accountability

A recurring concern was cognitive offloading: the erosion of instinct, skill and insight when tasks are handed to AI. One example was a vulnerability-identification tool that flags indicators before a human sees them. That may improve outcomes, but it can also weaken the professional’s own recognition and judgement. Accountability becomes fragile when people no longer do, see or fully understand the work.

That does not mean such tools have no place. It means firms must think carefully about how accountability is retained, how the task is being performed, and whether the tool helps the professional learn and add judgement rather than simply devolve the work. Used well, these systems can generate new data and insight that support better decisions. The key is that curiosity and judgement still sit with the professional.

Risk of miscategorisation (false positives and negatives)

Where AI is being implemented to help triage or support humans for example in supporting vulnerable customers, there will always be a risk of false identification. This risk exists of course with humans too, however additional consideration has to be given to the implications of systemic false positives and negatives.

A false positive in this scenario might lead to a customer being miscategorised as vulnerable, and given guidance and support tailored accordingly. If the firm has designed all products and services inclusively this may not lead to significant harm in itself, particularly if the engaged human in the loop is alert to and able to correct for this risk.

A false negative might miscategorise the customer as not vulnerable, and the engaged human in the loop may have narrow options and support considerations. This might have happened without AI, however there is an additional risk introduced here that the human has assumed that the AI is correct, and they are less alert to and able to adapt their guidance to respond. This is akin to the [‘too many quality controllers’ conundrum](#) wherein BMW discovered that after a certain point, more quality controllers led to more errors, since each assumed one of their colleagues must have spotted it before them.

NB vulnerability indicators inferred from AI should never be treated or recorded as factual until corroborated with the customer by a human. This is to respect the data accuracy principle in UK GDPR and mitigate against the risk of miscategorisation.

Augmented, not artificial, intelligence

Several participants preferred the term augmented intelligence. Much current practice already keeps a human in the chain, though rarely where the human sets out a mission, takes ownership of the result and has full visibility and control of all handoffs. There is a spectrum from conventional automation, wherein a human sets a series of defined operational tasks to be executed without autonomy, to agentic, wherein a human sets a brief and allows the agent to choose its own processes. There is a risk that augmentation can mislead users and stakeholders: people may not realise how much of what they see has already been shaped by earlier systems and biases. The 2008 financial crisis was cited as a warning about low-friction handoffs that hide complexity and weaken accountability. As AI use grows, firms will need a stronger focus on professional ethics. The audit profession’s emphasis on professional scepticism was seen as a useful model for insurance and personal finance.

The ethical challenge: from population fairness to the individual

Participants discussed a number of ethical challenges¹ to conventional AI policies. It was broadly agreed that conventional policies focus on fairness at the level of the population or “our customers as a group,” because that is how systems are optimised. A professional signs up to act in the best interests of each individual customer, and that difference forces a different set of questions. How could something that benefits most customers fail this individual; can I explain to this person how the decision was reached; and can I challenge my organisation where a whole group of customers loses out?

Suitability and the chatbot question

AI may produce an answer that is technically correct and statistically defensible yet professionally wrong for a customer’s circumstances. Customer service or guidance Chatbots illustrate the risk: a customer who says they need a product may simply have that view confirmed, where a trained professional would probe what they have not considered. Accountability (who is responsible when a chatbot effectively gives advice) came up repeatedly. The conclusion was that tools must have a clearly defined purpose, deliberately chosen and managed rather than assumed.

Fairness, power and the reclassification of risk

A major ethical question was what happens when stronger techniques and richer data reclassify a previously insurable risk as uninsurable, as in wildfire or flood modelling. AI-enabled hyper-personalisation raised concern about weakening the pooling principle at the heart of insurance. Participants disagreed, usefully, on whether more data necessarily excludes more people: it can also mean more accurate and granular pricing. The shared concern was power. AI does not create this dilemma, but makes it more frequent and more acute. Consumer Duty, foreseeable harm and public scrutiny already apply.

Conflicts of interest and the duty to challenge

Participants also stressed that some AI risks are conflicts-of-interest issues. Systems may optimise for retention, profit or claims-cost reduction over customer value. Businesses have always had to balance those interests, but rapid AI adoption can make the trade-offs harder to see. Professionals therefore need the knowledge, autonomy and psychological safety to ask whether a system is doing the job they remain accountable for, and to challenge or escalate where necessary. More broadly, the central question is not the governance of AI in isolation, but the governance of a firm, embodying its values and culture, as AI changes it.

¹ See appendix 2

Governance, procurement and operational resilience

A dedicated AI committee, or the existing risk function?

Participants debated whether AI needs its own governance committee. Several argued against another silo: the core risk process does not change and issues such as bias extend beyond AI, so the priority is the right expertise at each stage and clear accountabilities, including a single executive owner. Others said firms are creating dedicated AI committees because issues were bouncing around existing structures, often alongside cyber. Strong parallels were drawn with cybersecurity governance, including the case for AI scenario planning and simulation exercises.

Are existing governance and control frameworks still effective?

Most existing governance and control frameworks were designed for conventional systems, where processes behave as specified and controls can be tested against predictable outputs. AI that includes autonomous decision-making (e.g. systems that adapt, infer or choose their own path to an objective) challenge this. Firms should therefore review their existing frameworks to test whether they remain effective when such technologies are in use.

Vendor due diligence and third-party risk

Procurement was another concern. Questions a standard process asks of a conventional software supplier may miss AI-specific risks, and firms may assume controls that do not in fact apply. Issues such as data provenance, how decisions are reached, data residency, equivalency, and whether a vendor uses client data to improve its own products may not be covered. Even where they are, operational, risk and legal teams may struggle to assess them given the pace of change.

Participants pointed to work the CII could build on rather than recreate, including TechUK's material on AI procurement and assurance, the OECD's definitions, and an interactive tool for judging whether a vendor response is inadequate, acceptable or strong². A useful organising idea was to triage use cases: some can proceed within defined conditions, while others need governance or risk-committee approval, especially where sensitive or vulnerable-customer data is involved.

Review, unwinding and operational resilience

Participants argued for disciplined post-implementation review: did the deployment meet its objectives, has its purpose shifted, and if it is not delivering, can it be unwound safely? That includes understanding what capability would be lost if AI tooling were withdrawn. Cost also mattered. Compute and token costs make the value question unavoidable, and the clearest gains so far were seen in areas such as coding rather than core underwriting. Two further points were largely missing from current policies: the sustainability and energy intensity of AI, and the fact that this is a transformation programme, with historically high failure rates, in which investment in people and culture has lagged far behind investment in infrastructure.

² See appendix 3

~70%³

Digital transformation initiatives
failure to meet objectives

3-6%⁴

Proportion of total AI budgets spent
on employee training and AI literacy

Skills, capability and CPD

Participants agreed that training is one of the most important and most underestimated requirements of responsible AI. It must go well beyond “an hour on Copilot”. Capability-building needs to cover what AI is, where it is useful, where it is risky, and the ethical questions it raises. It should also feed back into the prior question of whether a function needs AI at all and what the business case is. Poor understanding creates real risk, from data protection breaches when staff paste client information into public models to the lazy assumption that AI always improves productivity.

A strong theme was that AI changes the competency profile of the professional. The routine early-career work through which judgement was once built is increasingly automated, so newer entrants may no longer see under the bonnet. Different generations need different training: experienced professionals may be able to test machine output against years of context, while newer entrants may treat it as authoritative. At the same time, digital natives may be more fluent with interfaces and coding and better able to direct agentic systems. The most important competence, participants suggested, is knowing where you are not competent and asking good questions. That points to CPD focused on behaviours such as critical thinking, professional scepticism, insight and challenge, assessed through realistic scenarios rather than treated as ‘soft skills’, with less importance than ‘hard’ technical knowledge. A skills-audit element was also proposed for governance, covering whether firms and committees have the right capabilities now and are maintaining them for the future.

³ Boston Consulting Group (2020), [Flipping the Odds of Digital Transformation Success](#), which found that 70% of digital transformations fall short of their objectives

⁴ Allocation frameworks aggregated May 15, 2026 by Presenc.ai from [Deloitte's State of AI in the Enterprise 2026 report](#), BCG's 10/20/70 framework documentation, [Tredence's 2026 AI spending analysis](#)

What the CII and other bodies can do

In the final part of the session, participants considered what the CII, the wider sector and other bodies could do next, building where possible on work that already exists.

Define what “good” looks like – as a framework, not boilerplate

The CII should set out what a good AI policy looks like for insurance and financial planning firms, aligned to the Code of Ethics and good customer outcomes. Participants were clear that this should be a flexible framework that helps firms express their own culture, strategy and operating model, not a template to copy. One strong example was a small financial planning firm’s policy that named the tools it used, explained why, set expectations for use and linked to its vulnerable-customer policy. A live question is how such policies stay current as tools change week by week.

Make corporate Chartered status a genuine value proposition

Because the CII can directly influence only firms seeking corporate Chartered status, participants encouraged a pull as well as push approach: make responsible AI something firms want because it creates value, not just another compliance demand. That could include a clear consumer-facing signal such as a badge or kitemark, showing that a firm can demonstrate responsible AI processes and outcomes. Guidance should also cover any beneficiary of a client, including commercial and business customers, not only retail consumers.

Develop AI fluency and behavioural CPD

There is an urgent short-term education need. Participants suggested AI should become part of professional training and recognised that many students already use it in learning, adapting materials and contributing to non-invigilated coursework. Passing off AI output as one’s own should still be treated as misconduct, but the ability to work with AI is now part of professional reality and should be addressed by adapting learning and assessment to test how young professionals augment their work. Given the pace of change and variety of use cases, CPD was seen as more useful than qualifications for building and maintaining capability, especially through case studies and scenarios that test judgement rather than recall. The Code of Ethics was also seen as a strong basis for AI CPD, giving professionals questions to ask and situations to think through rather than focusing only on tool use.

Provide benchmarking, maturity curves and live playbooks

Because AI use and capability are moving faster than regulation, firms would benefit from objective comparison with peers. Participants saw value in CII-led benchmarking on what a good policy looks like, AI maturity curves, and what is and is not working in practice, potentially with other professional bodies where AI spans functions such as insurance, HR or finance. Use cases and playbooks should be treated as living documents, with trade and professional bodies well placed to aggregate anonymised experience across firms and ‘raise the floor’ of practice.

Work across bodies on standards, governance and procurement

Participants called for sector-wide agreement on governance frameworks, including clearer triage for when AI use needs formal governance and how third-party risk should be managed, and for the CII to signpost credible external resources rather than duplicate them. Opportunities included linking broker standards work with CII good practice and Chartered CPD, collaborating with other Chartered bodies, and promoting better use of sandboxes. Across all of this, the CII has a role in speaking clearly from an ethical standpoint and keeping sight of the profession’s longer-term direction.

Conclusion and next steps

The discussion reflected both the pace of change and the profession's determination to meet it responsibly. The clearest conclusion was that responsible AI is less about inventing new rules than about applying enduring principles: duty to each individual customer, personal accountability, professional scepticism and the duty to challenge. Much of what firms need already exists in regulation and the Code of Ethics. The task is to help professionals apply judgement well and keep purpose, people and trust at the centre of adoption.

To take this forward, the CII will:

- develop practical outputs on AI policy frameworks, CPD, benchmarking, use cases and playbooks; and
- continue the conversation with participants and collaborators as this work develops.

Appendix 1: Artificial intelligence in insurance and personal finance - a policy scaffold

Artificial Intelligence in Insurance and Personal finance – a policy scaffold

This appendix outlines the areas an insurance or personal finance firm should consider in its AI governance and policy.

Chartered firms will be expected to publish a policy that reflects their organisational purpose and operating priorities and sets out their approach to the principles below. No single policy will suit every firm, and firms should not treat this scaffold as a checklist to adopt in full. Advice-led firms will face different opportunities, risks and priorities from firms focused mainly on product delivery. Policies will also vary by the size and complexity of the business and by whether it serves consumers, intermediaries or businesses.

NB This scaffold focuses on the use of AI within insurance and personal finance and is not intended to cover every governance issue raised by AI in a business. Many other resources already address broader good practice, such as intellectual property misuse and hallucinations or errors in AI-generated text.

1. Purpose and Scope

This policy sets out how the organisation develops, deploys, and governs Artificial Intelligence (AI) in a way that is safe, fair, transparent, and aligned with our obligations to customers, regulators, and society. It applies to all AI systems used in underwriting, pricing, claims, fraud detection, customer interactions, and internal operations, including third party AI tools.

2. Principles for Responsible AI

2.1 Fairness and Non Discrimination

- AI systems must not result in unfair or unlawful discrimination.
- All models undergo bias testing at design, deployment, and periodic review.
- All AI-enabled processes should undergo individual level ethical review, not just model level fairness.

2.2 Transparency and Explainability

- Customers must be able to understand when and how AI has influenced decisions that affect them. In practice may mean: disclosure at the point of decision, plain-language explanations of the AI and human roles and clear routes to request human review or challenge.
- High impact decisions require clear, human readable explanations.
- Reliance on extensive terms and conditions should be avoided in communicating the use of AI in preference for proportionate customer-centric language
- Black box models are prohibited in regulated decision making without approved justification.
- Communications should explain not just the technical factors that produced a decision, but the professional rationale applicable to the individual circumstance.

2.3 Human Oversight and Accountability

- High risk decisions (e.g., underwriting, claims repudiation) require active and engaged human in the loop.
- Clear ownership is assigned for every AI system, including model performance and outcomes.
- Staff must be trained to understand AI limitations and ethical principles and escalation routes, including whistleblowing where appropriate.
- AI outputs must be filtered through professional judgement before being applied to individual customers.

2.4 Safety, Robustness and Security

- AI systems must be resilient to errors, manipulation, and data drift.
- All models undergo stress testing, validation, and ongoing monitoring.
- Third party AI must meet equivalent standards.

2.5 Legal and Regulatory Compliance, Privacy and Data Protection

- The firm must ensure that its use of AI complies with all applicable laws and regulations – including, but not limited to, UK GDPR and wider data protection law, the Consumer Duty, equality legislation, and any AI-specific requirements as they emerge.
- Use of personal data must be proportionate and justified.
- Synthetic or anonymised data should be used where possible.

2.6 Ethical Boundaries (“Red Lines”)

The organisation will not use AI for:

- Social scoring or behavioural manipulation.
- Emotion recognition for eligibility or pricing.
- Fully automated decisions that significantly affect customers without human review.
- Any use that can reasonably be expected to undermine customer autonomy or trust.

3. Governance and Oversight

3.1 AI Governance Committee or subcommittee of existing oversight function

- Provides oversight, approves high risk models, and monitors compliance and ethical outcomes.
- Includes representatives from risk, compliance, actuarial, data science, and customer functions.
- Actively monitors and corrects for customer trust and perception of trust in outcomes enabled by AI.
- Actively seeks conflicts of interest between commercial and customer outcomes of AI enabled processes.

3.2 Model Lifecycle Management

- Standardised processes for design, validation, deployment, monitoring, and retirement.
- Mandatory documentation for all models, including training data, assumptions, and limitations.

3.3 Third Party and Vendor Controls

- All external AI tools undergo due diligence, risk assessment, and contractual controls.
- Vendors must provide transparency on data sources, model logic, and performance.
- Controls and training to avoid inappropriate staff use of unauthorised tools

4. Customer Protections

- Clear communication when AI influences outcomes.
- Accessible routes for human review and challenge.
- Regular audits to ensure outcomes remain fair and aligned with customer interests.

5. Reporting and Escalation

- Material AI risks must be escalated to the Executive Risk Committee (or appropriate subcommittee).
- Incidents involving unfair outcomes, data breaches, or model failures must be reported promptly.
- Annual reporting to the Board on AI usage, risks, and improvements.
- Whistleblowing procedures, training and empowerment for staff to challenge AI outputs.

6. Review Cycle

This policy is reviewed at least annually or proactively following significant regulatory, operational or technological developments.

Appendix 2: Ethical challenges

The CII Code of Ethics and the Digital Ethics Companion set expectations that may go beyond regulatory or commercial good practice in the interests of the public and the reputation of the profession. The examples below show where corporate AI policies could go further.

1. Integrity: Acting in the Best Interests of Each Individual Customer

Where the gap is

AI policies tend to focus on fairness and non discrimination at a population level. But the CII Code requires **acting in the best interests of each individual customer**.

The ethical challenge

AI systems optimise for groups, not individuals. This creates tensions such as:

- Pricing algorithms that are “fair” statistically but disadvantage a specific customer
- Claims automation that is efficient overall but fails to recognise individual vulnerability
- Fraud models that flag customers based on correlations rather than meaningful evidence

What’s missing

A requirement for individual level ethical review, not just model level fairness.

2. Competence & Care: Ensuring Advice and Decisions Are Appropriate for the Customer

Where the gap is

AI policies typically cover human oversight, but not the professional duty of **suitability**.

The ethical challenge

AI may recommend actions that are technically correct, statistically justified but **professionally inappropriate** for a specific customer’s circumstances

Examples:

- Automated underwriting that doesn’t recognise atypical but legitimate risk factors
- Chatbots giving generic guidance that is interpreted as advice
- AI driven nudges that influence customer decisions without ensuring suitability

What’s missing

A requirement that **AI outputs must be filtered through professional judgement** before being applied to individual customers.

3. Transparency: Explaining Not Just the Decision, but the Professional Rationale

Where the gap is

AI policies typically cover explainability, but the CII Code requires **clear, honest, and professionally meaningful explanations**.

The ethical challenge

AI explanations often fall into:

- technical descriptions (“the model weighted X and Y”)
- generic statements (“your risk profile changed”)

Neither meets the ethical requirement for clarity, relevance, honesty and professional justification.

What’s missing

A standard for **professionally adequate explanations**, not just technically adequate ones.

4. Professionalism: Avoiding Conduct That Damages Trust in the Profession

Where the gap is

AI policies typically cover governance, but not **reputational and societal trust impacts**.

The ethical challenge

Even ethically designed AI can undermine trust if:

- customers feel decisions are impersonal or opaque
- automation appears to prioritise efficiency over empathy
- vulnerable customers feel unsupported
- pricing feels “algorithmically opportunistic”

The CII Code requires professionals to **uphold the reputation of the profession**, which includes perceptions, not just outcomes.

What’s missing

A requirement to assess **customer trust impacts**, not just compliance and fairness.

5. Conflicts of Interest: AI Optimising for the Firm Over the Customer

Where the gap is

AI policies typically cover fairness but not **conflicts of interest**, which the CII Code explicitly prohibits.

The ethical challenge

AI systems may:

- optimise pricing for profit rather than customer value
- recommend products that maximise retention rather than suitability
- prioritise claims cost reduction over fairness

These are not “bias” issues — they are **conflict of interest issues**.

What’s missing

Explicit controls ensuring AI does not create or amplify conflicts of interest.

6. Duty to Challenge Unethical Practice

Where the gap is

AI policies typically cover governance, but not the **professional duty to challenge**.

The ethical challenge

If staff see AI producing:

- unfair outcomes
- misleading explanations
- inappropriate recommendations

The Ethical Code requires them to **challenge and escalate**. Policies don’t typically explicitly embed this cultural expectation.

What’s missing

A requirement for:

- whistleblowing routes for AI concerns
- psychological safety for challenging AI decisions
- training on when and how to challenge AI outputs

Appendix 3: Further resources

Participants referenced a number of resources the CII could build on or signpost.

Getting started and good practice

- [CII AI insurance and personal finance – a policy scaffold](#)
- [ABI – Practical ideas for getting started with responsible AI](#)
- [CII Code of Ethics and Digital Ethics Companion](#)
- [OECD AI Principles overview – guidance and definitions of AI](#)
- DSIT [Implementing the UK's AI Regulatory Principles](#)

Contextual references

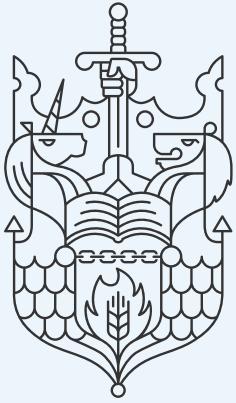
- [Harvard Business Review: The Psychological costs of adopting AI](#)
- [LFBF and CSFI Framing the risks facing financial services in the Gen AI era](#)
- [Insurance Times: The risks and rewards of hyper-personalised insurance models](#)

Procurement, assurance and governance

- [Tech UK Lessons from the AI procurement frontline: Beyond the contract and buying blind: Rethinking AI procurement as a governance function](#)
- Reframe Venture [due diligence tool](#)

Related CII material

- [CII managing customer vulnerability guidance](#) – relevant where AI is used to identify or support vulnerable customers.
- [Artificial Intelligence and vulnerable customers. Roundtable summary report.](#)



Chartered
Insurance
Institute

Responsible AI: from policy to practice Roundtable summary report