

DON'T OUTSOURCE TRUST TO AI – USE AI TO SCALE TRUSTWORTHY ADVICE

by Prashant Bharadwaj & Dominic Houlder



Prashant Bharadwaj advises global accounting, tax, legal and financial services sector firms on AI strategy and execution – helping them deploy AI responsibly while preserving the trust their clients and regulators expect. He previously held senior roles in private equity advisory and strategy consulting at Monitor Deloitte, and operational leadership at a Motorola JV. He holds a masters in computer science and an MBA from London Business School.



Dominic Houlder is an Adjunct Professor of Strategy at London Business School and a former BCG strategy consultant. He has deep experience guiding global professional services firms –including PwC, EY, and Deloitte – through strategic transformation. Dominic holds an MA from Cambridge University and an MBA from Stanford Graduate School of Business.

The next big scandal in professional and financial services may not begin with fraud or negligence. It may begin with a regulator asking a simpler, more damaging question: who, exactly, owned the AI output? The AI opportunity is real, but only if your clients and regulators still trust you on the other side. And no one is exempt from the trust question.

Professional and financial services leaders face a defining choice: deploy AI with trust at its centre, or react to events, outsourcing accountability and eroding credibility.

Trust is the bottleneck, not the technology

The commercial stakes are already high. Significant PE money has already gone into the global accounting and tax market¹. This won't stop at professional services. Wealth management, mortgage brokerages, retail banks, building societies and insurance all face the same forces. If AI can remake tax advisory, it can remake mortgage lending. Tech start-ups across the world are rapidly building AI systems for these regulated sectors.

Early adopters are showing how this plays out, in both directions.

Last year, Deloitte Australia submitted a federal government report containing fabricated citations: AI-generated errors that slipped through review. Deloitte issued a partial refund. The story became a global boardroom case study within days².

More recently in May 2026, Pinsent Masons was admonished by London's high court after its lawyers made inaccurate submissions to a judge based on AI-generated documents. Pinsent Masons referred itself to the SRA over AI-related failures³.

Allen & Overy deployed an AI contract tool to 3,500 lawyers across 43 offices⁴. But leadership made one thing non-negotiable: validate everything before it leaves the firm⁵. The difference was not the technology. It was whether the firm had designed its solution to foster trust before AI reached the client.

Anthropic's research⁶ found that AI could automate far more professional and financial services work in theory than firms are deploying in practice. That gap is not a capability problem. It is a trust problem. Most firms today cannot answer a simple question: when AI gets it wrong, and it will, whose career ends? Who will pay?

This is your dilemma: how do you adopt AI at pace without corroding the trust on which your firm, and your profession, relies?

Our view is straightforward, even if its implications are not:

Do not outsource trust to AI – use AI to scale trustworthy advice.



Trust must remain with people and the institutions that back them. AI's role is to help those people deliver better service, faster and more consistently – not to replace their judgement, ethics, or accountability.

What outsourcing trust looks like today

When firms slide into the outsourcing of trust, there is no dramatic handover to machines. It looks mundane.

A manager drops AI-generated paragraphs into a report designed for a pre-AI world. The review process is unchanged. There is no stage at which AI output is specifically challenged, no record of how key claims are supported.

Chatbots answer questions on allowances or compensation as if they were trained staff, but without constraints or supervision. When errors surface, regulators show little patience for blaming the system.

AI handles 30% of an employee's workload, but no one redirects that capacity. It vanishes into busywork. The opportunity to reinvest in complex client work or enhanced governance is completely squandered.

Much of today's regulatory and professional guidance is sensible in direction – "be safe, be fair, be accountable" – but it often misdiagnoses where harm enters. It assumes the main risk sits inside the algorithm, so it leans heavily on ideas like "explainability" and boilerplate statements.

Professional standards make a similar error. They correctly insist that humans remain responsible, but they rarely specify what good supervision looks like for an AI-enabled solution that is inherently probabilistic yet sounds authoritative at scale.

This is what outsourcing trust currently looks like: a gradual erosion of clear human ownership over what is said and signed.

The no-naked-AI principle

Here is the strategic shift that changes everything. Your AI may help draft the work. Your customers' and regulators' AI will increasingly check it.

This is no longer hypothetical. In August 2025, a UK tribunal ordered HMRC⁷ to disclose whether it had used AI in assessing R&D tax relief claims⁸. If HMRC cannot hide its AI use from scrutiny, neither can you.

In user testing of a tax planning AI agent, we observed that users instinctively validated outputs by running them through Microsoft Copilot, Grok, and ChatGPT. They were not being difficult. They were being rational. If one AI produces advice, another AI can check it.

In this new world, every deliverable becomes evidence. It must show where claims came from, what was checked, where AI was used, and where a human took responsibility.

Trust at scale doesn't come from explaining a black box; it comes from governing the process around it with evidence, accountability, and recourse built in. No AI-enabled advice should reach a client without a check to suit the professional discipline and a named human sign-off from someone who understands they own the output.

Simply put: no naked AI.

From agentic to accountable

Enterprises are racing towards agentic AI – systems that act autonomously with minimal human intervention. It may be Big Tech's ideal vision. It may not be yours. The competitive logic is compelling: go fully agentic, cut headcount, maximise productivity – until a single unowned error undermines the entire business.

A 98% accuracy rate may be excellent for a consumer AI chatbot. It is career-ending for a tax adviser who signs off 500 high-net-worth returns a year, and finds ten of those challenged by tax authorities. The adviser cannot say 'my error rate was within tolerance'.

No AI-enabled advice should reach a client without a check to suit the professional discipline and a named human sign-off from someone who understands they own the output. Simply put: no naked AI.

For high-trust sectors, where advice carries personal liability and regulatory scrutiny, autonomy cannot come at the expense of accountability. Autopilot has existed for decades, yet planes still have pilots.

Consider a practical example. A tax planning AI agent can run multiple what-if calculations in seconds for a high-net-worth individual. But the tax adviser still has to sense-check the output, put it in context, and turn it into advice the client can act on. With an AI doer – human checker design, the client gets faster, more consistent analysis. The adviser spends less time on spreadsheets and more time on judgement. Trust stays where it belongs.

Five trust-at-scale tests before you commit to AI

Before committing to a major AI transformation, put trust at the centre, not technology, and consider applying these five tests. Each test is designed for a world where your clients' or regulators' AI is watching.

- 1. Whose name is on this?** When AI is involved, can you say in one sentence who is responsible if the advice is wrong?
- 2. Can you show the trail?** Can you show, end-to-end, how AI-assisted work is produced, checked, logged, and signed off?
- 3. How will their AI interrogate yours?** Are you designing for a world where clients' and regulators' AI systems interrogate your outputs?
- 4. Where did the freed human time go?** Where does AI-freed capacity go? Into fewer people

doing the same things, or into explicitly defined value-adding activities?

- 5. How are the lessons being applied?** How do you handle AI errors and near-misses? Is there ExCom-level monitoring to detect, learn, and adjust?

The pattern is there

There is a long and expensive memory for failures of accountability. Boeing's MCAS system was automated software that overrode pilot judgement. Engineers flagged concerns, but executives trusted the system and its cost efficiencies over the humans raising alarms. 346 people died⁹. In the UK, the Post Office scandal destroyed lives. Over 900 sub-postmasters were wrongly prosecuted because executives trusted a faulty computer system over the humans who said it was wrong¹⁰. Blind faith in automation, inadequate oversight, and years before anyone noticed. Until somebody did.

In this new world, every deliverable becomes evidence. It must show where claims came from, what was checked, where AI was used, and where a human took responsibility.

Professional and financial services firms adopting AI at scale are building the same accountability gap. The automation is more sophisticated due to advances in technology, but the oversight challenge is identical.

Within the next three to five years, a major U.S. or U.K. financial institution or accounting or legal firm will face regulatory action, not for AI bias, but for failing to maintain human accountability over AI-enabled advice. The regulator will ask the question that ended every previous defence: but who owned the output?

The firms that put trust at the centre and pass the five tests will have an answer. The firms that don't will discover that the consequences are existential, not just reputational. 



REFERENCES

1. Craig Jourdan, Elizabeth Todd, "Unpacking Private Capital's Growing Interest in Professional Services: The Opportunities and Challenges," Ropes & Gray Insights, June 2024, <https://www.ropesgray.com/en/insights/viewpoints/102j851/unpacking-private-capitals-growing-interest-in-professional-services-the-opport>
2. Nino Paoli, "Deloitte Was Caught Using AI in a \$290,000 Report to Help the Australian Government Crack Down on Welfare After a Researcher Flagged Hallucinations," Fortune, October 7, 2025, <https://fortune.com/2025/10/07/deloitte-ai-australia-government-report-hallucinations-technology-290000-refund/>
3. John Hyde, "Law Firm Pinsent Masons and Three Solicitors Referred to SRA After 'Astonishing' AI Failures," May 2026, <https://www.lawgazette.co.uk/news/pinsents-refers-itself-to-sra-over-ai-failures/5126895.article>
4. A&O Shearman, "ContractMatrix," <https://www.aoshearman.com/en/expertise/markets-innovation-group/contractmatrix>
5. Chris Stokel-Walker, "Generative AI Is Coming for the Lawyers," Wired, February 21, 2023, <https://www.wired.com/story/chatgpt-generative-ai-is-coming-for-the-lawyers/>
6. Maxim Massenkoff and Peter McCrory, "Labor Market Impacts of AI: A New Measure and Early Evidence," Anthropic, March 5, 2026, <https://www.anthropic.com/research/labor-market-impacts>
7. HMRC: the United Kingdom's tax, payments, and customs authority
8. STEP, "Tax Court Orders HMRC to Reveal Use of AI to Assess Tax Relief Claims," STEP Industry News, August 28, 2025, <https://www.step.org/industry-news/tax-court-orders-hmrc-reveal-use-ai-assess-tax-relief-claims>
9. Bill George, "Why Boeing's Problems with the 737 MAX Began More Than 25 Years Ago," January 24, 2024, <https://www.library.hbs.edu/working-knowledge/why-boeings-problems-with-737-max-began-more-than-25-years-ago>
10. Karl Flinders, "Post Office Horizon Scandal Explained: Everything You Need to Know," Computer Weekly, March 19, 2026, <https://www.computerweekly.com/feature/Post-Office-Horizon-scandal-explained-everything-you-need-to-know>