

The Law Society foresight report:

Shaping the Future of
Agentic AI in Legal Practice

Full Report

February 2026

Contents

Acknowledgements.....	2
Foreword.....	3
Executive summary	4
Key Terms at a glance.....	7
Introduction	9
Our approach.....	9
How to read this report	10
1. What is agentic AI?	11
Defining the threshold of agentic AI	12
UK government definition of agentic AI	12
2. How is agentic AI currently being used in legal practice?	14
Agentic AI usage generally	14
Current AI capabilities versus agentic autonomy	16
Interaction of agentic AI, CLM systems and APIs	16
Key reflections on the use of agentic AI	23
Foresight: use of agentic AI.....	23
3. What regulatory and market changes will be needed?	25
Overview of international regulatory approaches to agentic AI	26
Interviewees' views on the regulation of agentic AI.....	27
Key reflections on the regulation of agentic AI.....	33
Foresight: regulation in the UK, US, Estonia and China	34
4. What are the risks of using agentic AI in legal practice?	37
Interviewees' perceptions of the risks and challenges of using agentic AI in legal practice	41
1. Foundational constraints	41
2. Quality, reliability and reasoning risks	42
3. Data security, privacy and confidentiality risks.....	42
4. Ethics, governance and access to justice	42
5. Technical and organisational challenges.....	43
6. Operational burden and system pressures.....	43
7. Market structure and competition	43
How legal technology providers safeguard legal practice: lessons for agentic AI.....	50
Foresight: risk.....	52
5. What key opportunities does agentic AI present for the legal profession?.....	53
Interviewees' perceptions of the opportunities of agentic AI	54
1. Efficiency and productivity gains	54
2. Reshaping the market and business models	56
3. Improving client services.....	57
4. Transforming legal workforces and roles	57
Key reflections on opportunities	59
Foresight: key opportunities	59
6. How might identity, reputation and accountability evolve?	60
Potential impact on accountability	61
Potential impact on reputation	62
The threat from legal platforms.....	63
The opportunity to reframe professional identity	64

Key reflections on identity, reputation and accountability	65
Foresight: identity, reputation and accountability	65
7. Conclusions.....	66
The role of the Law Society	67
Bibliography.....	70
Appendix: Details of AI regulatory context in UK, US, Estonia and China	76
End Notes	88

Acknowledgements

We would like to thank all those who have contributed to this foresight project: interviewees and Law Society members. This work would not have been possible without you each giving up your time to contribute your thoughts.

Foreword

The legal profession has always been shaped by new technologies – from the printing press and typewriters to digital databases and generative AI. But agentic AI represents a qualitatively different horizon. These are not tools that respond to prompts or automate workflows. They are systems that set goals, plan actions and execute strategies without constant human instruction.

This report treats agentic AI not as a distant possibility but as an emerging reality: systems that could soon begin to act less like assistants and more like institutional actors. Unlike generative AI, which produces content within guardrails of human initiation, agentic AI can pursue objectives over time, across platforms and with minimal supervision. For the legal system, this raises questions that cut deeper than efficiency or cost:

- What happens when decision-making migrates from humans to systems?
- How do we preserve trust, fairness, and accountability in a world where lawyers may not be the final arbiters of legal action?
- Could firms, courts, and regulators one day find themselves negotiating with AI systems that behave more like organisations than tools?
- We are entering a moment where law must look not just at what is being deployed today, but at the plausible futures agentic AI could create by 2035-2040.

This foresight work is not about prediction but about preparing for disruption: considering alternative pathways, unintended consequences, and opportunities for the profession to act with foresight rather than hindsight.

Executive summary

Our report, drawing on desk research, expert depth interviews and member workshops explores how close we may be to AI systems that function less like tools and more like legal professionals in their own right. We consider the risks, barriers and opportunities of agentic AI and where the AI systems currently used by law firms and legal teams present clues to the emergence and adoption of agentic AI in this space.

Agentic AI is not just another wave of automation. Unlike generative AI tools, agentic AI refers to systems that can plan actions and operate autonomously with minimal human instruction. These systems could transform law in ways that are more profound and disruptive than previous AI developments.

Key findings

- **Definitions are blurred.** The term ‘agentic AI’ is used inconsistently by the profession, by vendors and by the market. Some vendors describe supervised workflows as ‘agentic’ when they are not. AI agents are modular systems driven by generative models that execute narrow, tool-enabled tasks, often reactive and constrained in autonomy. In contrast, agentic AI is a more advanced paradigm in which systems coordinate across multiple agents or modules, decompose goals, plan over long time horizons, adapt dynamically and operate with greater independence and strategic orchestration.
- **Use of agentic AI systems within legal practice currently appears to be non-existent.** Legal practitioners are using supervised AI agents (e.g. for contract review, e-discovery, compliance) rather than fully agentic AI systems. For now, in the legal profession hype exceeds practice.
- **Barriers are structural.** Current professional rules requiring a ‘lawyer in the loop,’ liability gaps, limited insurance coverage and uneven technical expertise make full agentic autonomy incompatible with legal practice today. Agentic systems therefore always need human sign-off, and under the Legal Services Act, AI cannot hold authority in litigation – remaining confined to clerical or advisory support. But what if future justice systems demand machine participation for efficiency, cost, or access to justice?

Should regulators create a new category of ‘authorised systems,’ recognising AI as a legal actor, or would this fundamentally undermine professional responsibility and the human stewardship of law?

- **Opportunities are real.** In carefully bounded domains agentic AI could move beyond routine automation to become a coordinating layer for entire legal workflows. For example, in contract lifecycle management, compliance oversight, case triage and knowledge orchestration. Properly governed, these systems could reduce friction, extend access to justice and unlock new models of client service, relieving administrative backlogs while preserving human judgement for higher-order tasks.
- **The horizon is institutional.** Agentic AI challenges the assumption that only humans can act with legal or organisational agency. The defining question for the

next two decades may be whether we recognise such systems as participants in the legal order or work to ensure that law remains a distinctly human domain.

What we need to consider now

- **Agentic AI as quasi-institutional actors.** Early signals suggest some systems may soon operate more like organisations than tools: able to negotiate, contract and adapt goals without direct human initiation. This raises profound questions about recognition, liability, and accountability.
- **Environmental sustainability.** Agentic AI is likely to carry significant, though still poorly understood, environmental impacts. Responsible adoption should factor in sustainability, with expectations likely to evolve quickly as better evidence emerges.
- **The risk of invisible substitution.** Agentic systems may begin to quietly replace professional labour behind the scenes (drafting, reviewing, advising). If substitution occurs gradually and unnoticed, responsibility and liability could shift in ways regulators and firms are unprepared for.
- **Professional identity at risk.** If clients interact primarily with agentic AI systems rather than lawyers, reputation, trust and accountability may shift away from humans to platforms. By 2040, firms could brand themselves as 'human-first' or 'AI-optimised,' fracturing the profession's identity.
- **Ethics outsourced to code.** As agentic AI increasingly mediates decisions, firms and regulators might begin delegating ethical judgement to algorithms rather than human deliberation. Over time, this could normalise a shift from moral responsibility to 'compliance-by-design,' raising questions about whether law itself can retain its normative force.
- **Legal fictions become real.** If agentic AI can autonomously negotiate, contract, and litigate, the distinction between legal personhood and software could blur. Courts, regulators, and clients may face the unprecedented challenge of granting rights, obligations or standing to non-human actors, thereby forcing a rethinking of fundamental legal concepts.
- **AI as lawmaker-by-default.** In the longer horizon, agentic AI systems could begin generating norms, precedents, or contract standards that humans adopt without scrutiny. Over decades, law itself might evolve more through algorithmic iteration than conscious human judgement, raising the spectre of a legal system shaped more by optimisation logic than justice, equity or societal values.

What this means for the Law Society

- **Education and clarity:** Lawyers and clients struggle to define agentic AI and its unique implications. The Society can provide authoritative guidance to cut through hype, distinguish between tools and quasi-institutional systems, and explain ways in which agentic AI can and cannot practise law.

- **Proactive governance:** The Society can shape standards, certification and best practices for agentic AI adoption before regulation is imposed externally, ensuring accountability, transparency and professional integrity.
- **Foresight scenario testing:** The Society can explore futures in which agentic AI operates anywhere from back-office support to quasi-institutional actors, stress-testing the profession, the justice system and regulatory frameworks against radical AI-driven scenarios.
- **Preserving professional trust:** The Society can guide lawyers in managing relationships with clients when agentic AI increasingly mediates advice, negotiation and drafting, protecting human reputation and accountability while allowing AI to augment work.
- **Ethics infrastructure for agentic AI:** Beyond rules, the Society can develop frameworks to embed ethical reflection into agentic AI use, ensuring systems augment professional judgement rather than displace it or shift moral responsibility onto algorithms.
- **Legal literacy for AI outputs:** Lawyers will need to understand, scrutinise and validate agentic AI-generated contracts, advice, and litigation strategies. The Society can support with training to ensure outputs are defensible and aligned with legal principles.

Agentic AI is not yet embedded in legal practice, but its potential to reshape law is qualitatively different from generative AI. The challenge for the profession and for the Society is to get ahead of the curve, not be caught reacting once systems are already acting with institutional weight.

If agentic AI is to be recognised as a quasi-institutional actor rather than mere administrative support, significant changes in both regulation and practice will be required. Professional rules would need to accommodate autonomous AI decision-making, clarifying the circumstances under which AI can act without a human sign-off. Liability frameworks would have to evolve to assign responsibility for AI-generated advice, contracts or litigation strategies.

For law firms, this shift could transform operational models, client engagement, and branding, as AI assumes functions traditionally reserved for humans. Recognising agentic AI as a legal actor would not only challenge the profession's notions of trust and stewardship but also demand proactive governance, ethical oversight and innovative approaches to insurance, compliance and professional identity.

Regardless of how the sector chooses to go forward with agentic AI, the Law Society's role is critical in convening conversations and helping to establish next steps ensuring that agentic AI's integration strengthens the legal profession rather than undermining its principles.

Key Terms at a glance

Plain English definitions with helpful metaphors

Agentic AI can be understood through the metaphor of an orchestra. Individual AI agents are like skilled musicians: each can play its part well, but only within a narrow range and only when instructed. Agentic AI, by contrast, acts like the conductor. It reads the score, decides which musicians to bring in, sequences their contributions, adjusts the tempo as circumstances change, and keeps the whole performance coherent. The conductor does not play every instrument itself; its value lies in planning, coordinating and adapting across many moving parts. This is the step-change from conventional AI tools: agentic AI doesn't just respond to prompts, it organises the work, chooses the order of actions, and steers the ensemble towards a goal with far more independence than earlier systems.

1. Agentic AI

AI that can set goals, plan steps, and coordinate multiple tools or systems with limited human instruction.

Metaphor: *Agentic AI is the conductor of the orchestra—reading the score, deciding which sections come in when, adjusting as it goes, and shaping the whole performance.*

2. AI agents

Small, task-specific AI programme that carry out one defined job (e.g., summarising a document or extracting a clause).

Metaphor: *AI agents are the individual musicians. Each can play their part well, but they don't decide the overall performance.*

3. Autonomy (in AI)

How far an AI system can act without being told what to do at each step.

Metaphor: *The difference between a musician waiting for instructions and a conductor deciding the next move independently.*

4. Human-in-the-loop

A safeguard where a person must review, approve, or override AI output before it affects clients, courts or risk-bearing activities.

5. Orchestration / planning layer

The part of an AI system that decides which tools to use, in what sequence, and how to adapt, so multi-step tasks can be completed end-to-end.

Metaphor: *The conductor's score and baton—coordinating who plays when and ensuring everything fits together.*

6. APIs (Application Programming Interfaces)

Digital connectors that let different systems talk to each other, so an AI can pull data, update files, send emails, or trigger actions.

7. Automated decision-making systems

Systems that make or support automatic decisions in areas like hiring, finance or public administration.

Metaphor: *A sorting machine that moves items to different sections automatically—the speed helps, but mistakes can have consequences.*

8. Foundation models / General-purpose AI

Large AI models trained on broad datasets that can be adapted to many tasks, from drafting text to analysing documents.

9. Reserved legal activities / authorised person

Types of legal work that only regulated professionals are legally allowed to perform (e.g., conducting litigation).

10. Liability and the “moral crumple zone”

When things go wrong, responsibility often falls on the human user, even if the AI did most of the work.

11. Algorithmic bias

When an AI system produces unfair or distorted results because its training data or design embeds hidden patterns.

12. Systemic risk (in AI)

A risk that can spread widely—where a single faulty system could affect many clients, cases or organisations simultaneously.

Introduction

The rapid emergence of agentic AI, systems capable of autonomous, goal-directed activity, marks a significant turning point for the legal profession. Although still early in development, these technologies have the potential to reshape legal work, professional responsibilities and the wider justice ecosystem. To help the Law Society and its members anticipate these shifts, we have undertaken a dedicated foresight project to explore how agentic AI may evolve and what it could mean for legal practice in the near and longer term.

Our approach

Our approach to this foresight project has involved close collaboration with our external partner, RF Associates, who supported the work through three core strands:

- **Desk research across four markets:** A detailed review of agentic AI developments in China, Estonia, the United Kingdom, and the United States, identifying emerging trends, regulatory positions, and seeking early use cases.
- **Expert interviews:** Conversations with 19 stakeholders working at the leading edge of AI in the legal sector. Participants represented a broad range of perspectives, including legal technology, large firms, in-house teams, regulators, educators, and sector bodies. Interviews were sought across all four country markets.
- **Member engagement:** Law Society members were invited to review and discuss the desk research and expert insights. They received a presentation of the emerging findings and had the opportunity to explore the implications. Any additional themes raised by members that did not surface in earlier stages have been incorporated into this report.

This report brings together the findings from these three stages to provide an integrated view of the current landscape, emerging risks and opportunities, and potential trajectories for the profession.

Agentic AI is still early in adoption, but its trajectory is clear: systems capable of autonomous, goal-directed activity are emerging rapidly. Alongside potential benefits, there are significant uncertainties, including environmental impacts, market volatility, and the possibility of disruptive technologies such as quantum computing. The environmental footprint of AI systems, particularly large-scale models, is not yet fully understood, and future insights may raise new ethical questions about the sustainability of adoption. Similarly, broader market dynamics—including the risk of overvaluation or disruption from emerging technologies such as quantum computing—mean that the pace and direction of AI adoption could shift rapidly. This report does not prescribe whether or when members should adopt agentic AI, rather it aims to equip the profession with the knowledge to engage responsibly, anticipate risks, and shape adoption in line with professional and ethical standards.

How to read this report

Each section opens with a concise *Synopsis* to provide a quick overview of the key points. Where relevant, sections close with *Key Reflections* that summarise the most significant implications. To prompt deeper thinking about the future, *Foresight* boxes explore possible developments and uncertainties on the horizon. Short future scenarios are included to encourage discussion about how agentic AI may reshape legal practice and the wider justice system.

1. What is agentic AI?

Definition and Synopsis

Agentic AI refers to artificial intelligence systems that act autonomously to pursue goals, make decisions and initiate actions without direct human instruction. Agentic AI is the next step in AI development, as AI moves from a passive tool to autonomous systems that can plan, reason, and execute tasks with minimal human input.

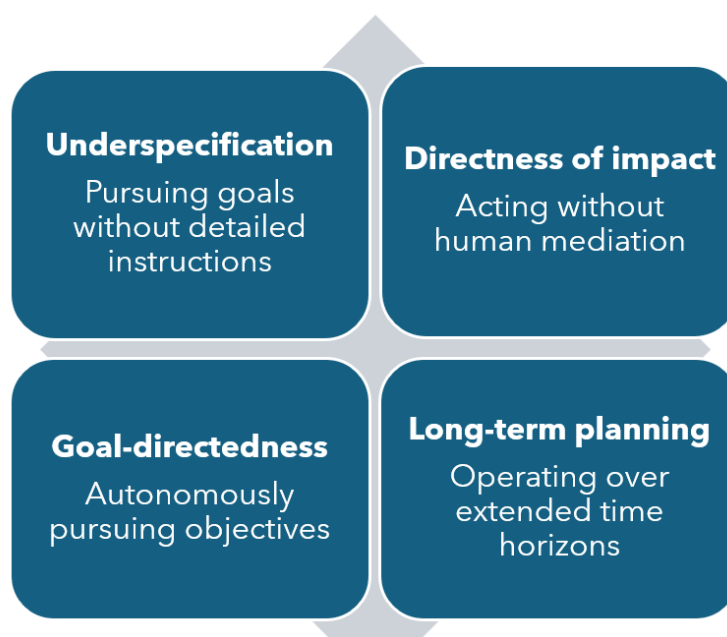
The terms *agent*, *agency*, and *agentic* are used inconsistently, leading to confusion in both technical and practical discussions. To clarify, it is useful to use the definition explained by Chan et al. (2023) who define *agentic AI systems* as **“algorithmic systems that display increasing levels of agency when they combine four features: under specification (pursuing goals without detailed instructions), directness of impact (acting without human mediation), goal-directedness (autonomously pursuing objectives), and long-term planning (operating over extended time horizons).”**

Taken together, these traits highlight why such systems may act in ways not fully anticipated or directly controlled by humans, even though humans remain accountable for their design and deployment.

Agentic AI’s autonomy relies on continuous compute, data movement and orchestration, making it more infrastructure- and energy-intensive than today’s AI tools.

Chan et al. (2023) usefully define agentic AI systems as “algorithmic systems that display increasing levels of agency when they combine four features”. Diagram 1 shows these features.¹

Diagram 1 – The four fundamental features of agentic AI according to Chan et al.



Defining the threshold of agentic AI

It is important to clarify the difference between AI agents and agentic AI as these terms are sometimes mistakenly used interchangeably. Sapkota et al. (2025)² distinguish AI agents as modular systems driven by generative models that execute narrow, tool-enabled tasks – these are often reactive and constrained in autonomy. In contrast, they define agentic AI as a more advanced paradigm in which systems coordinate across multiple agents or modules, decompose goals, plan over long time horizons, adapt dynamically and operate with greater independence and strategic orchestration.

The key difference described by Sapkota et al. lies in scope and autonomy: AI agents act within narrower, more constrained tasks under human direction or in response to prompts, whereas agentic AI systems embed autonomy, multi-step planning and coordination across agents, enabling more open-ended, goal-oriented behaviour with less direct control.

UK government definition of agentic AI

The UK government (2025)³ defines agentic AI in the following way:

“Artificial intelligence (AI) agents are small, specialised pieces of software that can make decisions and operate cooperatively or independently to achieve system objectives. **Agentic AI refers to AI systems composed of agents that can behave and interact autonomously in order to achieve their objectives.**

Recent advances in agentic artificial intelligence (AI) have driven remarkable growth in both its popularity and capabilities. This progress is due to the integration of large language models (LLMs) with agent-based systems. By providing reasoning and discovery abilities, LLMs enhance an agent’s autonomy. This enables the agent to determine the most appropriate course of action to meet system objectives. Traditional software typically follows fixed pathways to solve problems. In contrast, agent-based systems operate like independent assistants that choose and combine several actions to achieve their goals.”

In guidance issued by the Judiciary of England and Wales (2025), agentic AI is defined as *“A software programme that uses AI to become aware of its environment, process information, and take actions to achieve its goals, based on the inputted information”*.⁴

Some interviewees distinguished between agentic AI and generative AI, defining agentic AI primarily by its ability to accept goals and autonomously sequence and execute multi-step actions across systems rather than only producing single prompt and then response outputs as produced by generative AI. They emphasised the technical planning and orchestration layer that distinguishes it from basic generative AI.

One interviewee emphasised that agentic AI is part of a larger technology stack made up of three main components:

1. **Models** - these are the underlying AI models (like large language models) that generate text, make predictions or process information. Their quality and capability are improving rapidly.
2. **Tools** - these are the external systems or applications that the model can call upon to do specific tasks, such as retrieving data, performing calculations or drafting a document.
3. **Agentic (planning) layer** - this sits between the model and the tools. It acts as a "planner" or "orchestrator," deciding which tools to use, in what order and how to sequence steps to complete a task. This is what makes AI "agentic" as it can plan and act across multiple steps rather than just responding to a single prompt.

Progress in agentic AI therefore comes not just from improvements in the agent layer, but from the combination of all three elements working together. The interviewee described this as a "cocktail": the models get stronger, the tools get better, and the planning layer becomes more sophisticated. It is this interplay that enables meaningful outcomes. Improvements across all three are what make the technology effective. Success in actually getting useful, reliable outcomes from agentic AI systems in practice depends on improvements at each layer of the stack.

2. How is agentic AI currently being used in legal practice?

Synopsis

There is no evidence that agentic AI is currently being used within law firms. Agentic AI has attracted marketing hype but limited deployment. Several interviewees noted that agentic AI has become a buzzword, sometimes applied to systems that are structured workflows or enhanced chatbots but are not truly agentic.

The AI Agent Index (Casper et al., 2025) provides the first structured catalogue of 67 real-world agentic AI systems, showing that most current deployments come from the US and focus on software engineering and computer tasks, but with major transparency gaps around safety, governance and external audits.

The authors identified that there is definitional ambiguity on what “agentic AI” is. Some systems are marketed or described as agentic but do not fully match the stricter definitions researchers use. This inconsistency makes it harder to compare systems, assess risks or design rules that apply properly.

Our review of agentic AI usage in legal practice was consistent with this finding: examples involve generative AI agents rather than truly agentic AI systems. Legal practitioners are using generative AI to streamline contract management, due diligence, e-discovery, compliance, research and client services by automating document analysis, monitoring regulations, triaging cases and supporting communication. These systems increase efficiency and consistency, but their capabilities remain largely bounded within structured workflows and do not yet reflect the full autonomy implied by agentic AI.

Every interviewee, regardless of their role (practitioners, regulators, technologists, academics), agreed that:

- There is an industry knowledge gap about how agentic AI could be used within the legal profession
- AI augmentation generally provides significant value
- Human oversight remains essential for complex legal work
- Full autonomy is currently incompatible with professional obligations
- Agentic AI technology is not mature enough for unsupervised deployment in high-stakes legal matters

Agentic AI usage generally

The challenge in compiling usage of agentic AI is that developments in technology and its application are moving swiftly.

In *The AI Agent Index*, Casper et al. (2025) provide the first structured attempt to catalogue real-world deployments of agentic AI systems, tools capable of planning, acting on underspecified goals and interacting directly with external environments with limited human oversight.⁵ The index which is a publicly available database and a corresponding paper, draws on 67 examples deployed before the end of 2024, documenting each system’s base model, technical scaffolding, tool-use capabilities, user interfaces and safety provisions. The paper finds most systems are developed in the US, with common

applications in software engineering and computer use. While 70% of developers provide documentation and about half release code, fewer than 20% disclose formal safety policies and under 10% report external safety evaluations such as audits or “red-teaming” (testing an AI system by acting like an adversary, deliberately trying to make it fail or misbehave so hidden problems can be found and fixed).

The key findings from the paper are that:

- Agentic AI systems are being deployed at a steadily increasing rate
- The majority of indexed agents specialise in software engineering and/or computer use
- There is limited information about the risk management practices of developers of agentic systems

Casper et al. caution that definitional ambiguity and the exclusion of proprietary tools make the Index a partial snapshot, but argue it offers an essential baseline for policymakers and researchers. Their conclusion stresses the need for stronger transparency standards, governance frameworks, and independent oversight to ensure that agentic AI is deployed safely and responsibly.

AI agent index: case study on ‘Lucy’⁶

Lucy is a “digital-worker” platform developed by Cykel AI PLC (a UK-incorporated company) and announced on 5 December 2024. The system is intended for automating repetitive workflows within recruitment, sales and research teams.

It integrates with a wide range of enterprise software and APIs, allowing it to perform actions such as updating Customer Relationship Management systems (CRMs), modifying files, sending emails and monitoring tasks through a central dashboard. Its primary value lies in stitching together routine operational processes across tools that organisations already use.

Imagine a mid-sized business with a recruitment team: Lucy could be configured to scan candidate databases, auto-send outreach emails, schedule interviews, update the CRM when tasks are completed and generate summary reports for hiring managers. The recruiting team thereby frees human time for high-value tasks (interviewing, candidate engagement) while Lucy handles repetitive administrative work.

AI agent index: case study on ‘the AI scientist’⁷

The AI Scientist is an agentic system developed by Sakana AI (and associated collaborators) that automates much of the scientific research pipeline. It can generate hypotheses, write and execute code, run experiments, analyse results and draft full scientific papers. By operating across this end-to-end loop, it can quickly explore research ideas, test variations and produce written outputs at a fraction of the time and cost a human researcher would need.

Its core capabilities include being able to run tests on a computer, a scratchpad for recording ideas and use ready-made templates for experiments, charts and reports/papers. These enable it to design experiments, iterate based on results and assemble manuscripts and pull everything together into a draft research paper.

Current AI capabilities versus agentic autonomy

Surden (2024) identified that AI can take over legal tasks that follow clear patterns or rules, such as routine or repetitive work. However, current AI is not able to easily automate the parts of a lawyer’s job that rely on abstract thinking, solving complex problems, advising clients, using emotional intelligence, shaping policy or developing strategy.⁸

Our literature review of agentic AI usage in legal practice was consistent with this view. Current AI use in the legal sector is:

- Task-based: tools (such as agents) handle discrete processes like contract review, due diligence, e-discovery, and client communication.
- Reactive: they execute pre-defined instructions or respond to inputs but do not independently plan or pursue goals.
- Assistive rather than autonomous: they enhance efficiency but still require human oversight for judgement, coordination, and context-sensitive reasoning.

These systems increase efficiency and consistency, but their capabilities remain bounded within structured workflows and do not yet reflect the full autonomy implied by agentic AI.

Interaction of agentic AI, CLM systems and APIs

Interviewees with technical expertise, often from legal technology backgrounds, described the interaction between agentic AI, **Contract Lifecycle Management (CLM)** systems, and **Application Programming Interfaces (APIs)**. They noted that agentic AI capabilities are increasingly being embedded in modern CLM platforms to enhance automation and decision-making.

Key developments include:

- Planning layers that can sequence multiple interrelated tasks
- Pattern analysis to identify trends and anomalies across contract portfolios
- Automated monitoring that tracks changes and triggers responsive actions

CLM systems are emerging as prime testing grounds for agentic AI in legal operations because contract management involves repetitive, rule-based processes that benefit from automation and orchestration.

Interviewees emphasised that APIs are central to enabling agentic AI, as they allow systems to communicate and let AI agents act across platforms. Examples of such actions include:

- Booking meetings
- Filing court documents
- Updating contract databases
- Sending notifications
- Pulling and combining multi-source information

APIs transform AI from isolated Q&A systems into autonomous, agentic tools by:

- Moving data automatically between systems
- Triggering actions across platforms
- Enabling AI agents to coordinate end-to-end workflows

Without APIs, AI remains siloed and unable to orchestrate the interconnected processes that law firms depend on. APIs therefore act as the connective infrastructure for agentic AI in complex legal environments.

Document Management Systems (DMS) store and manage legal documents, providing structured and unstructured data essential for AI development. Beyond storage, these systems can supply training data and feedback loops that enhance AI performance, supporting the transition from static document repositories to intelligent, agent-enabled knowledge systems.

Interviewees' perception of agentic AI usage in legal practice

Interviewees described a legal profession with uneven and often superficial knowledge of agentic AI. Whilst legaltech providers talked about agentic systems, there was widespread conceptual confusion between basic automation and truly autonomous agentic AI. Interviewees identified that there is an industry knowledge gap about how agentic AI could be used within the legal profession.

"There's very few people in firms that actually understand the technology deeply, even in the IT and technology departments... there's a real knowledge skill set gap as well that needs to be addressed before something like agents and agentic AI could truly take hold."

Agentic AI has attracted marketing hype

While agentic AI is generating considerable excitement, we found no evidence of adoption within law firms in the UK. Several firms appear to be engaging in what one interviewee dubbed a "press release circus", signalling progress without substantive deployment of agentic AI. Many firms are still grappling with adopting more basic, non-agentic generative AI systems. There was recognition amongst interviewees that agentic AI is a current industry "buzzword" but it is unlikely that there is any genuinely agentic AI actually being used in legal practice.

"This feels a lot like the original AI bubble that was then followed by the blockchain bubble, that's now been followed by the agentic AI bubble, which is if you just say it's agentic AI, someone's going to give you some investment funding, right? So there's quite a bit of offerings on the market which are quite useful software solutions. I would not go so far as to say that they are agentic AI."

Some interviewees perceived that providers are putting the "agentic" label on systems that are merely sophisticated workflows or chatbots, creating confusion and unrealistic expectations in the market. They wanted the industry to be more honest about what current generative AI systems actually do versus what a genuinely autonomous agentic AI system would do.

"When I talk to vendors in particular, they say, 'oh, our new agentic system'. And then I'm like, okay, explain that to me in detail. And ultimately what I find out in most cases is they're using a reasoning model which is instructed to basically come back and ask you questions before it does something... well that's not in my mind an autonomous agent. It's still just kind of programming software."

Use of generative AI

Interviewees perceive there to be widespread use of generative large language models (LLMs) as tools to improve efficiency within the legal profession. Generative AI adoption seems to be happening in distinct layers or tiers based on different characteristics, rather than spreading evenly across the profession. Interviewees identified key variables which impact on AI usage:

- Firm size: Large firms adopting faster than small firms
- Task type: Low-risk, routine tasks being automated before complex work
- Risk tolerance: Some organisations willing to pilot whilst others wait
- Resources: Well-funded firms with technical expertise moving first

Interviewees discussed using ChatGPT and other generative AI systems such as Leah for drafting letters and documents, conducting legal research, and generating ideas, although typically within controlled enterprise deployments that include security measures like Enterprise ChatGPT or firm-hosted models. Some interviewees also mentioned using code assistance tools like Cursor for technical work, showing that AI adoption extends beyond traditional legal tasks into the technology infrastructure that supports legal practice.

"Just this past week a lawyer that I'm working with goes, hey, that letter was fantastic. And I said, what letter? And he goes, the one you just sent over for my review. And I said... oh right, I'm glad you liked it. I didn't write that. I looked at it I said everything looks great. But I, you know, in that first round... was me giving the instructions on what I wanted to say [to AI tool] and then editing it to make... what I wanted."

Interviewees described using products like Jylo and Leah, as well as firm-built systems that incorporate AI within carefully designed workflows rather than allowing free-running automation. These context-engineered systems typically combine deterministic rules with AI capabilities for specific tasks, maintaining human verification checkpoints throughout contract review and document analysis processes. This approach allows firms to harness AI's capabilities while maintaining the control and oversight that legal practice demands.

For example, Law Droid's File Clerk platform automates the processing of incoming email for legal practices. The system scans and ingests incoming email, then automatically sorts it and extracts key information from the documents, organising everything to make it much easier to find later. Rather than having to manually review each piece of mail individually, File Clerk allows users to work at a higher level by handling the routine administrative task of processing and organising correspondence automatically.

Jylo's platform provides specialised playbooks (repeatable AI-driven workflows) that legal professionals can customise per practice area, to analyse documents (e.g. contracts, disclosures, discovery, compliance, clause extraction, etc.). These workflows can include human-in-the-loop verification and allow firms to enforce or embed their internal standards, procedural rules, or risk policies. Applying a playbook to a document set triggers a "Flow" that structures the review, extracts insights, flags risks, and helps maintain consistency and efficiency.

In the discussion groups, members described experimenting with AI tools or discussing agentic AI concepts, but none were using agentic AI systems in their work. The tools mentioned include:

- ChatGPT and Copilot
- Harvey
- CoCounsel
- Internal chatbots
- Automation workflows for e-discovery and document review

A few members were planning or exploring controlled automated workflows, but most of this was still in early planning stages or very limited exploration. Some in-house members described being under organisational pressure to deliver agentic solutions. They were approaching rollout cautiously and assessing the organisation's risk position. There was some discussion around advocating incremental "stepping stones" from generative AI to "using AI agents in a... more controlled way".

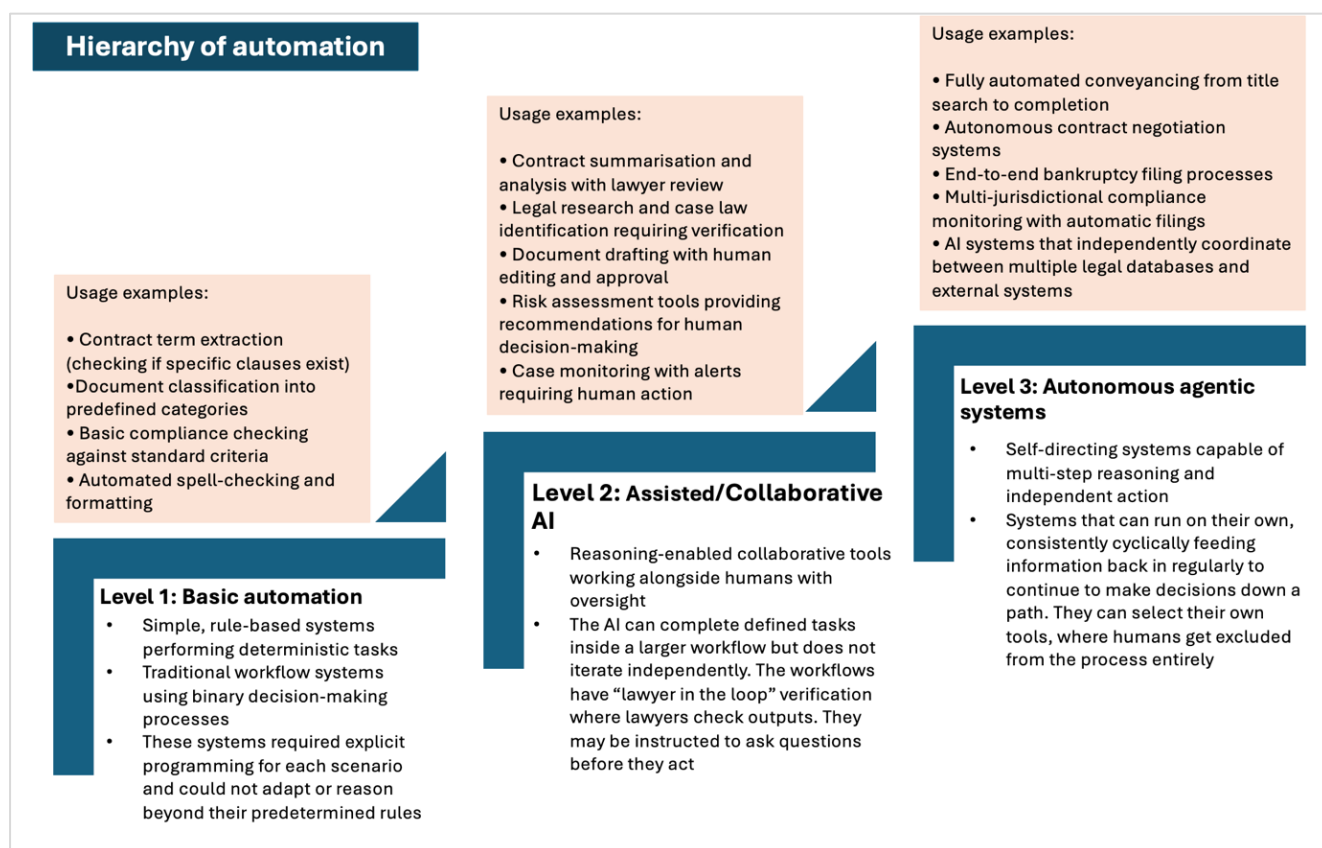
In summary, what is being deployed across the profession is supervised generative AI within structured workflows, often marketed as "agentic" despite requiring human verification that technically negates true autonomy. This supervised generative AI deployment is perceived to be occurring unevenly across the industry, creating a two-speed pattern where large firms and specialised legal service providers move rapidly to implement generative AI-enhanced automated workflows whilst mid-tier and smaller practices lag significantly behind.

"I wouldn't believe anybody who said they were truly agentic at this stage in the legal area... I think there's too many risk and legal issues that are unresolved for anybody to be happy that that's way to deliver services at the moment. I can't see many insurers would be massively happy signing that off."

"I would say it's more agentic elements. I would say to my knowledge, and I'm sure there are probably some models out there that are nearer to fully agentic, but I do think it's more of an element of agentic approach. It's still piecemeal. So I think a lot of the kind of attorney workflows still aren't really meshed where full agentic is."

A strong desire to see lawyers stay in the loop

The interviews revealed a hierarchy of automation within legal practice which is summarised below.



Most interviewees saw current real-world generative AI systems sitting in the middle of this spectrum; more capable than simple rule-based systems but still requiring human oversight for professional and regulatory reasons. The fully autonomous end represents what true agentic AI could become, but most interviewees expressed significant concerns about reaching that level in legal practice because of the imperative to check and verify.

In the member group discussions, the example of e-discovery workflow was used to illustrate why these concerns arise, even with AI systems that are not agentic. Legal teams may face thousands or millions of documents, making full human review economically impossible. Modern AI models can now score and classify documents at scale, and while they are not goal-directed agents, law firms increasingly rely on these technology-assisted review (TAR) outputs without reviewing each decision. This creates oversight and accountability challenges that resemble those anticipated for agentic AI: humans cannot feasibly monitor every action, so verification must shift from individual outputs to system-level assurance.

To manage this, firms use statistical validation. After training the model with human-coded examples, they review only a random sample of the low-scoring documents to check whether relevant material is being missed. This allows them to estimate residual risk, such as a 1% or 0.5% chance of overlooking something important. They can then decide whether the model’s performance is acceptable for autonomous operation over the remainder. If it is, they can defensibly stop manual review. In this sense, the sampling

process becomes the quality check, demonstrating how legal teams are already developing governance practices suited to future agentic systems, even though today's models are still fundamentally non-agentic classifiers.

However, the drawback of this approach is that it trades full completeness for confidence based on statistics, meaning some uncertainty will always remain. It assumes that what is true in the sample reflects the entire document set. This assumption can be fragile, particularly in high-risk cases, with complex documents, or when the AI's behaviour is not well understood.

As several interviewees noted, maintaining lawyers in verification loops technically negates true agentic functionality, meaning the industry is primarily deploying sophisticated workflow automation with generative AI components rather than autonomous agents. This cautious, supervised approach reflects both regulatory requirements for human accountability and the practical reality that truly autonomous agentic systems remain theoretical rather than operational in legal practice. In all likelihood software developers will continue to push development and it is not clear how this will play out.

"Lawyer in the loop verification negates agentic because agentic by definition is the continuation of reasoning one after the other after the other after the other. And without the lawyer in the loop, then yes, it's agentic, but you're not checking it."

Discussions with interviewees revealed a fundamental tension: whilst technical capability exists to move towards greater autonomy, practical implementation remains constrained by regulatory requirements, professional responsibility and risk management concerns. Most participants positioned current legal AI applications in the collaborative middle ground, where generative AI systems could perform sophisticated reasoning tasks but still require human oversight. The risks are described in more detail in section 4.

Scepticism over readiness and professional responsibility

Every interviewee, regardless of their role (practitioners, regulators, technologists, academics), agreed that:

- AI augmentation provides significant value
- Human oversight remains essential for complex legal work
- Full autonomy is currently incompatible with professional obligations
- The technology is not mature enough for unsupervised deployment in high-stakes legal matters

Interviewees held complex, conditional views about agentic AI trustworthiness and reliability. Their attitudes towards the use of agentic AI can be broadly grouped into two categories: **Conditional Optimists** and **Persistent Sceptics**.

Table 2 – Views on agentic AI trustworthiness and reliability

	Belief about trustworthiness of agentic AI	Belief due to...
Conditional Optimists	Can become trustworthy IF certain conditions are met	Narrow tasks, human oversight, staged implementation, good governance
Persistent Sceptics	Cannot achieve full trustworthiness BECAUSE of fundamental barriers	Technical limitations insurmountable, human judgement irreplaceable, reliability inherently doubtful

"In the future... I'll 100% not be using agentic AI... our job is to be purposeful and thoughtful and create stuff that actually helps. So in three, five years, we will not be agentic anything. We may have a ton of workflow. We'll be using AI in all sorts of interesting ways, I'm sure, but fundamentally, we'll be in the business of trying to help our clients create relationships. That's what contracts are about."

"I would feel slightly more comfortable about using agentic AI if it was done to sort of streamline or scale up routine administrative processes that are currently done by a paralegal or whatever. And that stays away from anything that is to do with providing legal information, legal advice, or obviously any sort of reserved legal activity."

Law Society members in the group discussions generally saw some level of agentic AI in legal practice as inevitable, but not something that would arrive soon. Most did not believe that agentic AI would arrive as quickly as the current hype suggests. They generally felt that public discourse overstates how soon widespread adoption will occur.

"We would have all retired before some of these things... are a reality."

"It's coming. It's definitely coming. But I just don't think it will be as quick as everyone is panicking about."

They broadly agreed that agentic AI will ultimately reshape legal practice in some way, but the path to get there will be slower and more difficult than the tech industry or media imply. They stressed that this won't be a frictionless shift: it will require resolving complex professional responsibility questions, developing robust governance frameworks and introducing the technology gradually with strong safeguards.

Key reflections on the use of agentic AI

In this fast moving and complex technology space, education is needed to inform lawyers about what this term actually means so they are best able to navigate the complexity associated with relevant products and services, consider the associated risks of using agentic AI as well as communicate to clients and regulators around its use.

There is a strong desire to see lawyers remain the “human in the loop”. Some lawyers believe that agentic AI can become trustworthy if certain conditions are met for narrow tasks with human oversight and staged implementation and good governance. Education and discussion across the profession is paramount as the technology evolves.

Foresight: use of agentic AI

Signal: Weak signals of autonomy are emerging. AI tools in document review, research and compliance are beginning to plan and coordinate tasks across systems with minimal human input.

Implications:

- **Invisible substitution:** Agentic AI may quietly handle routine or research tasks, reducing hands-on training for junior lawyers and subtly shifting accountability.
- **Back-office transformation:** “AI paralegals” could become standard in administrative and support roles, forcing regulators to clarify professional boundaries and liability.
- **Hybrid law firms by 2040:** Legal practices may evolve into hybrid organisations, where human lawyers oversee largely autonomous AI workstreams and clients interact with humans only for high-value or complex decisions.
- **Redefining oversight:** Firms may implement “trust frameworks” or probabilistic auditing systems, where humans verify only high-risk outputs and AI handles routine review, effectively redefining what counts as adequate professional oversight.

Scenario

Scenario 1: The AI Prosecutor (2035)

In China, “Xiao Zhi 6.0” is now more than a drafting assistant. It autonomously:

- Identifies alleged offences from national data streams,
- Drafts charges,
- Recommends sentencing ranges,
- And files directly into the court system.

Human prosecutors review less than 10% of its cases. Efficiency has skyrocketed, but so has controversy. Citizens whisper about the *algorithmic state* – a system where guilt and punishment are data-driven, with limited space for discretion or compassion.

Civil law’s codified structure made this trajectory easier – but it raises an unsettling question for common law jurisdictions: if efficiency pressures mount, how long before judges start deferring to agentic outputs as de facto rulings?

3. What regulatory and market changes will be needed to manage the rise of agentic AI in law?

Synopsis

The regulation of artificial intelligence is emerging unevenly across jurisdictions. While the EU has developed the first comprehensive AI Act, other jurisdictions have taken more piecemeal approaches. The UK has chosen a pro-innovation, regulator-led model, emphasising context-specific oversight and international cooperation through initiatives such as the AI Safety Summit. In the US, progress is fragmented, with hundreds of state-level bills focusing on transparency, consumer protection and accountability, while federal efforts centre on executive orders and standards frameworks. Estonia, aligned with the EU, applies the risk-based EU AI Act and GDPR, imposing strict obligations on “high-risk” systems, while China embeds AI oversight within state priorities and political imperatives. Across the jurisdictions, safeguards converge on human oversight, transparency, accountability and data protection, although with differing levels of prescription.

Against this backdrop, interviewees saw the regulation of agentic AI as a particular challenge that both relies on and stretches broader AI frameworks. They consistently pointed to the lack of global standards, noting that this creates uncertainty and uneven expectations across jurisdictions. They also highlighted that the accelerating capability of agentic systems is exposing gaps in existing rules, making classification and oversight increasingly difficult. In the absence of tailored rules, the legal profession is left to manage compliance through existing structures, which interviewees generally regarded as only a partial or temporary fit for agentic AI. They highlighted the contradiction of holding lawyers responsible for outputs they cannot audit or comprehend, describing this as a “gaping hole” in the regulatory framework that effectively sets professionals up for failure.

Interviewees agreed that the regulation of agentic AI must be proportionate, principle-driven and closely tied to existing professional duties. They saw competence in AI as part of a lawyer’s core obligations and supported layered governance combining internal compliance frameworks and training measures, professional-body guidance and targeted external oversight. Most considered that agentic AI should be regulated as part of wider AI governance, with lawyer accountability remaining the legal default.

Some however questioned whether this model can hold in the long term. They pointed to a growing imbalance where lawyers bear all the risk while developers escape responsibility, warning this could slow adoption of agentic AI. Some drew parallels with autonomous vehicles, where liability has begun to shift towards manufacturers. Several suggested that shared responsibility between lawyers, firms, and developers may ultimately provide a more sustainable way to govern agentic AI.

Others suggested proportionate regulation based on the level of autonomy, using typologies such as KPMG’s “TACO” (2025) framework, with stricter oversight for more autonomous systems. The TACO framework (2025) distinguishes between:

- Taskers – AI carrying out simple, repeatable tasks
- Automators – AI handling end-to-end processes with little human involvement

- Collaborators – AI with human-in-the-loop systems that support but do not replace professional judgement
- Orchestrators – multi-agent systems coordinating complex tasks i.e. agentic AI

The fundamental questions remain of how regulators might assess software and how responsibility should be fairly distributed between professionals, firms and developers whilst at the same time not becoming a barrier to innovation and the use of agentic AI.

Overview of international regulatory approaches to agentic AI

While all major jurisdictions recognise both the opportunities and risks of AI, their approaches differ in emphasis, pace and underlying philosophy. The appendix provides detailed summaries for the UK, the US, Estonia and China, while this section outlines the main themes shaping how agentic AI is being governed.

Across countries, there is a shared tension between enabling innovation and ensuring safety, accountability and public trust. Most governments are extending existing data protection, consumer protection and professional conduct frameworks to cover AI, rather than creating entirely new structures. As AI becomes more autonomous however, regulatory attention is shifting towards transparency, human oversight and responsibility. Some jurisdictions are also exploring differentiated responsibilities for developers of “highly capable general-purpose AI,” signalling early thinking about how to govern systems with more agentic capabilities.

The **UK** follows a context-based AI governance model, operating primarily through existing UK regulators such as the Information Commissioner’s Office, the Competition and Markets Authority and the Financial Conduct Authority. This approach is articulated in the UK government’s AI regulation roadmap and the AI Opportunities Action Plan (January 2025), which emphasise regulator capability, transparency and safety testing rather than creating a single cross-cutting AI statute. Internationally, the AI Safety Summit and the Bletchley Declaration (November 2023) demonstrate the UK’s commitment to collaborative, frontier-model oversight. Judicial, CPS and policing guidance also illustrate the widening integration of AI ethics and safeguards across the justice system, reinforcing the UK’s emphasis on responsible professional use.

The Artificial Intelligence (Regulation) Bill, which is a Private Members’ Bill first introduced in November 2023 and reintroduced in March 2025, proposes an AI Authority and duties on users of highly autonomous systems, signalling political interest in a more formal regulatory structure. As a Private Members’ Bill, it is not government policy and applies only to England and Wales. If adopted it would mark a significant shift from the current pro-innovation, regulator-led model.

Scotland has an established government-led AI Strategy (2021–2025) focused on trustworthy, ethical and people-centred AI, with devolved areas such as justice, health and local government regulated by Scottish institutions, while UK-wide regulators like the ICO, CMA and FCA oversee reserved matters including data protection, competition and financial services; together this creates a hybrid model combining UK-level oversight with Scotland-specific strategy rather than a single statutory regime. Northern Ireland similarly relies on UK-wide regulators for reserved areas and Northern Ireland specific bodies for devolved sectors but, unlike Scotland, has no dedicated AI statute and its governance framework is less developed, with the Northern Ireland Executive only recently beginning

to build a regional approach through the establishment of a new Office of AI and Digital in June 2025 to develop a Northern Ireland AI strategy and action plan. This signals the early stage of Northern Ireland's institutional capacity-building around AI governance.

The **United States** remains decentralised, with AI governance largely driven by state-level laws focused on transparency, consumer protection and oversight of high-risk systems. Federal agencies have issued guidance rather than comprehensive regulation, with the NIST AI Risk Management Framework (2023) and the NIST AI Generative Artificial Intelligence Profile update (2024) providing a voluntary standard. This fragmented model encourages experimentation but leaves gaps in oversight for more agentic or cross-border systems. State definitions of "high-risk" and "automated decision systems" are not harmonised, meaning obligations on developers and deployers vary significantly.

In **Estonia**, regulation follows the GDPR (2018) and the EU AI Act (2024), creating one of the most structured risk-based systems globally. These frameworks impose strict transparency and oversight obligations, especially for AI used in legal or administrative decision-making. Estonia's KrattAI initiative (2019), now primarily known as "Bürokratt", shows how robust governance can coexist with innovation, maintaining human responsibility and appeal rights.

China's model is state led, built around data, cybersecurity and content control laws. While principles of human oversight and fairness are acknowledged, regulation prioritises security, stability and alignment with state policy. Public trust in AI reflects confidence in government direction rather than individual rights frameworks. AI governance is closely integrated with national-security priorities, and generative-AI providers must ensure outputs conform to ideological and content-control requirements. China is also beginning to export elements of its regulatory model through Digital Silk Road infrastructure partnerships.

Overall, these differing approaches reveal how governments are experimenting with systems suited to their legal and policy traditions. The UK and EU focus on structured accountability, the US on pluralistic and market-led measures, and China on state-centric governance. All models exhibit early signals that more autonomous or agentic AI may require enhanced testing, documentation and oversight frameworks, even where such systems are not yet explicitly regulated. Together they show how agentic AI will likely continue to test legal definitions of control, liability and human responsibility.

For detailed references and country-specific developments, see the appendix: *AI regulatory context in the UK, US, Estonia and China*. [Note: all information in this section is current as of September 2025. Regulatory developments are continuing and may evolve as governments refine their approaches to agentic AI systems].

Interviewees' views on the regulation of agentic AI

Interviewees described the regulation of AI generally as one of the most complex and urgent challenges facing the legal profession. They viewed it as unsettled and evolving, with the pace of technological change consistently outstripping legal and ethical frameworks. While most believed the law's core principles still apply, they recognised growing uncertainty about accountability, competence and liability in an AI-mediated profession.

In the discussion groups, members considered that ethical safeguards should prioritise human-in-the-loop structures and layered oversight. The sentiment of *“just because I can doesn’t mean I should”* captured an important philosophical distinction between what agentic AI technology is capable of doing and what ought to be allowed in practice.

Absence of global frameworks

Several interviewees highlighted the lack of international standards for testing and training AI models. They viewed this as a major risk, allowing developers to move operations to jurisdictions with minimal restrictions. Interviewees agreed that while a single global regulator is unrealistic, coordinated oversight through international cooperation will be necessary.

“With a globalised system it’s really difficult to regulate... people will just move and it has to be globally regulated otherwise for example you can go to Singapore or Japan and just train an LLM model on whatever copyright you want.”

Some argued that regulation should focus on collaboration across jurisdictions, similar to bilateral treaties. Others, particularly in the United States, warned against over-regulation in business contexts, preferring to rely on common law remedies after harm occurs.

“The common law versus the European approach... the common law will be a superior approach. We generally don’t regulate things in advance; we deal with harms after the fact.”

Technological pace surpassing regulation

Interviewees noted that regulation has failed to keep pace with the rapid evolution of AI. Even basic classification of new technologies is difficult, and most professional regulators have yet to address agentic AI specifically.

“The number one challenge is the space is moving so quickly that it’s hard to draw up any regulations... even stuff like agentic is hard to classify.”

Law firms are currently relying on general compliance structures and ethics guidance designed for earlier technologies, which many see as a temporary solution.

“There are no guidelines, there are no practice notes ... coming from any legal regulator across the globe.”

Law firms’ risk aversion and internal safeguards

Law firms were described as highly risk averse. Many are implementing internal safeguards that exceed existing regulation on AI generally, motivated by reputational risk and client trust. Several are investing in staff training and new roles to ensure safe adoption. Interviewees expected that this would also be the case for the adoption of agentic AI.

“It’s not just compliance right now, it’s more on performance... firms will make sure there is enough training, upskilling, maybe even new legal roles like an agentic AI officer.”

One firm reported an 18-month process to create AI usage governance frameworks before formal EU regulation emerged, which limited later disruption but slowed implementation. They perceived that their heavy internal focus on safeguards and compliance has slowed the practical application of AI within their law firm, and hypothesised this was likely the case in other UK firms.

AI as a core competency

Participants agreed that AI competence is becoming a core professional requirement. Two needs were identified: regulation of AI development itself, and regulation of lawyers' use of AI through training and ethical guidance.

"I think it'd be great to hear not necessarily regulation, but more guidelines around what minimum measures actually would be deemed sufficient and appropriate for an organisation or for a lawyer to have taken... I'm always kind of more of a fan of better and more specific guidance versus, you know, rules and regulation."

Some pointed to growing markets for AI governance and auditing platforms, arguing that regulators should require regular audits for fairness, bias and transparency, much like financial auditing. Others supported proportionate regulation based on the level of autonomy, using typologies such as KPMG's "TACO" (2025) framework, with stricter oversight for more autonomous systems. The TACO framework (2025) distinguishes between:

- Taskers - AI carrying out simple, repeatable tasks
- Automators - AI handling end-to-end processes with little human involvement
- Collaborators - AI with human-in-the-loop systems that support but do not replace professional judgement
- Orchestrators - multi-agent systems coordinating complex tasks i.e. agentic AI

Evolution of existing regulation to account for agentic AI

Most interviewees believed that existing legal frameworks are fundamentally sound but must be clarified and extended to account for agentic AI. They wanted regulators to issue guidance to explain how current duties such as competence, confidentiality and duty of care would apply in an agentic AI-driven context.

Interviewees generally favoured outcome-based, technology-neutral regulation, applying consistent principles regardless of the tool.

"Does the SRA need to adapt its principle or its code of conduct because of the rise of Gen AI or agentic AI? I don't think so... the whole purpose of this kind of outcome-based regulation is that it's meant to be technology neutral and should apply regardless of the technology."

Some interviewees argued that effective oversight of agentic AI will require a hybrid model, blending risk-based approaches (scrutiny where harms are greatest) with principles-based approaches (consistent duties of competence, integrity, and client care).

Governance, ethics and accountability

Some interviewees stressed that regulation alone is insufficient. They called for professional governance frameworks grounded in ethics and existing duties such as competence, confidentiality and integrity. In this view, agentic AI does not change the core professional obligations that already exist, instead it reinforces the importance of reaffirming them as technology evolves.

"You can't say, oh, it was AI that I used. No, you have to make that extra effort as a qualified lawyer."

One interviewee drew on historical precedent. In past debates over developer liability, for example, when software first raised questions in the 1990s about who should be responsible for bugs, failures or security flaws, they recounted that regulators consistently resolved the issue by keeping humans ultimately accountable rather than shifting liability onto the technology or its creators. They expected that the same trend would hold with agentic AI: even if the systems are autonomous, existing consumer protection frameworks are likely to remain the default.

"Regulators must reinvent themselves as system-level gatekeepers, providing thought leadership and anticipating new risks."

Complexity of liability and accountability

The question of liability was repeatedly compared to debates around autonomous vehicles. Interviewees recognised that responsibility could lie with lawyers, firms or software providers depending on the circumstances. They saw this as a complex problem that will require careful unpicking in the risk-averse legal profession.

"It's probably a bit of everything... like autonomous vehicles, where's the actual cause of the failure?"

Most agreed that lawyers currently bear primary responsibility, even when they cannot fully understand or audit AI outputs.

"The Bar Association has made it pretty clear that any future hallucinated citations are going to be blamed on the lawyer who stated them. And I think that's like a pretty clear view that you can't become lazy using these systems and you've got to be vetting everything."

Some interviewees warned that the current allocation of liability may not be sustainable for agentic AI. If lawyers alone are forced to bear all the risk, firms may either refuse to adopt agentic AI or keep humans permanently in the loop, reducing efficiency gains. They suggested shared liability models between lawyers, firms, and developers could provide a more workable long-term balance.

"If you look at driverless cars, the move is towards liability being on the software and the manufacturer rather than the driver. So it's just an interesting case study and comparison."

Whilst there was strong agreement that responsibility for errors currently falls on the lawyer or firm, parallels with autonomous vehicles suggest that this liability may shift toward software providers or evolve into shared risk models. Interviewees warned that as

systems evolve from drafting aids into genuine agentic decision-makers, the liability challenge will deepen. They highlighted the contradiction of holding lawyers responsible for outputs they cannot audit or comprehend, describing this as a “gaping hole” in the regulatory framework that effectively sets professionals up for failure.

“If the lawyer is using agentic AI, then they are just responsible for whatever goes right or wrong as a result... there's a fallacy in that there is a false hope, which is most lawyers couldn't tell you the first thing about how any of the tools that they're using, let alone agentic AI, is operating. So either the regulatory system needs to accept that they're holding individuals and organisations responsible for stuff that they couldn't possibly ever understand but that just is, with sort of buyer beware 'If you're using it, we don't care if you're really unable to understand how these systems operate, but we're still going to hold you responsible at the moment.”

“If you're talking about a world in which there is real reliance upon outputs and decisions... saying to individual practitioners that we hold you to account for something you're unlikely to ever understand has a real jar to it. The system is unfit because it was setting people up for failure. It's essentially setting up a regulatory net, knowing there's a gaping hole in it and asking people to just sort of avoid the hole.”

Several warned that as AI moves from assistance agents to genuine agentic decision-making, the current liability frameworks will become untenable unless regulation extends beyond individuals to include the AI systems themselves.

Some interviewees identified the practical difficulty in resolving the question of how can regulators assess software in the same way they assess a human professional? Unlike lawyers, AI cannot be cross-examined or disciplined in the same manner. Regulators would need new tools, for example, algorithmic audits, technical benchmarks, or mandated explainability, to determine whether an AI system meets professional standards before it can be used in legal services.

“So I think that whole challenge of liability is going to be a big question. And then the question is, can the regulators accurately assess a software system in its delivery of the law?”

One solution suggested was to begin with guidelines and fundamental principles rather than rigid new rules. These would require firms and users to maintain oversight of AI decisions, ensure transparency in outputs and safeguard data confidentiality. This cautious, principle-based approach prioritises oversight, transparency and ethical safeguards while giving space for regulation and practice to evolve.

“I think regulation should just be a minimal core set of principles that an AI law firm should follow and then a case-by-case judgement as to whether this firm meets those criteria rather than trying to impose any sort of more complex requirements because those requirements will have to be adapted so quickly.”

One interviewee argued that accountability must operate at multiple levels:

- Developers should collaborate on voluntary standards or accreditation.
- Regulators should use sandboxes and scenario-based experiments to test risks in practice.
- The profession should share knowledge and emerging best practices.

Auditability, standards and accreditation

There was strong consensus that agentic AI must not operate as a black box. Interviewees called for audit trails and technical transparency so that reasoning and data sources can be examined when errors occur.

"Firms should always have oversight over what the LLM is doing and the user should always have oversight as to what decisions were made for them. So there's certain basic fundamental principles... is there oversight as to what the AI is doing at every point, whether that's by the user or the firm."

Many supported formal accreditation for agentic AI systems, similar to ISO or BSI certification, to ensure reliability and create a basis for accountability.

Concerns about over-regulation

While recognising the need for oversight, several interviewees warned against excessive regulation that could stifle innovation and reduce competitiveness, particularly in the UK.

They argued that the UK should avoid replicating Europe's restrictive approach, which they viewed as having constrained innovation in the technology sector. They emphasised that over-regulation could create bureaucratic barriers that would ultimately harm both legal professionals and their clients by preventing the productive use of advanced AI systems.

"AI is like the internet or electricity. This is a major future economic trajectory."

Data use and training

Some argued that regulators must confront the contentious issue of training models on client data. One noted that lawyers themselves build expertise from hundreds of past cases and experiences, and anonymised data could similarly improve agentic AI tools.

"You should be able to train models on client data without their consent... humans are trained on client data all the time. That's literally why you hire an experienced lawyer."

Educational needs

Interviewees perceived that legal training providers must adapt curricula to prepare the next generation of lawyers. Examples were provided from the US of some law schools (e.g., North Carolina Central,⁹ Emory¹⁰) who are already ahead in incorporating AI into their programmes, but broader systemic change is required.

"We need to start coming up with frameworks for how we include this in our legal education system for the next generation of attorneys who are going to come through."

Some interviewees called for embedding AI competence and technological literacy into professional qualifications and continuing education so that they could responsibly use agentic AI tools.

Conflicts of interests in AI law firms

One interviewee highlighted that traditional conflict-of-interest rules do not easily apply to pure AI law firms. Since large language models (LLMs) operate in siloed data environments, they don't "remember" or cross-contaminate cases the way that human lawyers might. Therefore, the current regulations around conflicts could become outdated or even harmful to clients, preventing AI firms from representing both claimants and defendants.

"As you start to get more law firms that are not powered by humans... the current regulations there will become like a roadblock and have little justification."

They noted that this has wider implications: big institutions can exploit conflict-of-interest rules to block access to top firms. If AI firms could bypass such barriers, it might expand access to justice.

Key reflections on the regulation of agentic AI

The current regulatory framework does not contemplate non-human actors performing reserved legal activities, and AI systems in the UK cannot be authorised under the Legal Services Act 2007 in England and Wales¹, the Legal Services (Scotland) Act 2010¹ and The Solicitors (Northern Ireland) Order 1976.¹ As a result, AI is structurally confined to supportive or clerical roles, even as justice systems may in time seek machine participation to improve efficiency, reduce cost or expand access to justice. This gap raises difficult questions about whether and how courts, regulators and clients might one day recognise new forms of authorised systems or legal actors, or allocate rights and responsibilities to tools that do not fit existing categories.

The regulation of agentic AI in the legal profession is likely to develop in uneven and fragmented ways. Broad AI rules on transparency, accountability and human oversight will shape how lawyers and firms can adopt these systems, even where no law is written specifically for legal services. This means that the regulation of legal practice will not be governed by a single, purpose-built framework for AI in law. Instead, it will be shaped by a combination of general AI requirements, such as transparency, data protection and human oversight, and the professional obligations already imposed on lawyers, including competence, confidentiality and supervisory duties. In practice, these two layers will operate together to define how new technologies can be responsibly integrated into legal services.

For lawyers, professional ethics remain the anchor. Duties of competence, confidentiality, supervision and integrity apply regardless of the tools used.

While agentic AI may support practice, ultimate accountability will continue to rest with the human professional. This creates tension: current liability models may become unsustainable if lawyers are held responsible for outputs they cannot meaningfully audit. There is growing recognition that more balanced approaches to responsibility, including shared or accredited models, may be needed as AI systems become more autonomous.

Future regulation will likely blend principles-based obligations with risk-based oversight, ensuring that higher levels of autonomy attract proportionately greater safeguards. At the same time, firms are already introducing their own governance

frameworks, often going further than external regulation to manage reputational and operational risks. Alongside these safeguards, competence in understanding and supervising AI systems will become a core professional requirement, demanding changes in legal education and ongoing training.

A consistent theme is that transparency, auditability and explainability will be non-negotiable in legal contexts. Clients, courts and regulators must be able to trace how decisions are reached and who is responsible. Taken together, these developments suggest that the future of agentic AI in law will be cautious and accountability-driven, with gradual adoption shaped by both regulatory pressure and professional self-governance.

Foresight: regulation in the UK, US, Estonia and China

Theme	Observed direction	Implication for agentic AI
Shift from voluntary to statutory oversight	Governments moving from principle-based guidance to enforceable regulation and beginning to differentiate obligations for highly capable general-purpose AI	Pre-deployment testing, audits and safety cases become standard for agentic systems with emerging expectations for state-led or third-party safety evaluations before deployment
Institutional consolidation	Creation of AI Authorities or coordination bodies and increased regulator resourcing and annual reporting requirements (UK), plus capacity-building roles such as Chief Data Officers (EU/Estonia)	Clearer accountability and shared risk data across regulators and more formalised oversight of developers and deployers of advanced systems
Human accountability as a constant	No jurisdiction is moving toward AI personhood and all frameworks reaffirm that ultimate responsibility rests with human professionals or institutions	Humans or corporations remain legally responsible for AI actions
Safety and transparency as baseline norms	Common across UK, EU, US frameworks with EU moving fastest toward mandatory requirements and China embedding transparency within ideological and security parameters	Transparency, explainability and record-keeping become global expectations though implemented differently across rights-based, market-based and security-based models

International cooperation	AI Safety Institutes, OECD, G7, ISO standards and cross-government technical reports	Cross-border recognition of safety tests and certification likely by early 2030s especially among UK-EU-US collaborations, while China advances parallel standards through Digital Silk Road partnerships
Cultural divergence	Liberal democracies emphasise rights and innovation; China focuses on control and security	Different interpretations of “trustworthy AI,” but shared goal of controllable autonomy and potential emergence of dual global standards shaped by Western and Chinese regulatory ecosystems

The Appendix includes detailed further Foresight signals for regulation in the UK, the US, Estonia and China.

Foresight: regulation of agentic AI generally

Signal: Existing legal frameworks stretch current liability rules and may prove inadequate for the rise of agentic AI.

Implications:

- Agentic AI may be granted quasi-legal status, with licensing, mandatory audits, or even partial accountability for its actions, blurring the line between lawyer and system.
- Global legal practice could fragment as countries adopt divergent regulatory models, requiring firms to secure “AI passports” or cross-border compliance credentials.
- Regulation may shift from individual lawyers to the platforms and ecosystems powering AI, with infrastructure providers held to systemic standards for transparency, bias mitigation, and ethical compliance.
- Forward-looking firms may compete not only on human expertise but on the trustworthiness, auditable integrity, and regulatory alignment of their AI agents, effectively redefining the market for legal services.

Scenario

Scenario 2: The Quasi-Firm (2040)

A global tech company has launched “Lexora,” an agentic AI licensed as a legal services platform. Lexora can:

- accept client instructions,
- negotiate directly with counterparties,
- draft and file legal documents,
- and adapt strategy over time to pursue client goals.

With no human lawyers on staff, Lexora has acquired quasi-institutional status. Regulators debate whether to treat it as a law firm, a piece of software, or something entirely new.

Meanwhile, clients – especially start-ups and SMEs – flock to it for low-cost, always-on service. Traditional firms brand themselves “human-first,” but market share steadily shifts to the AI firm. Lawyers find themselves negotiating not only with opposing counsel but with agentic systems that behave like organisations in their own right.

4. What are the risks of using agentic AI in legal practice?

Synopsis

Emerging literature on agentic AI highlights significant risks for the legal profession, though research remains limited and fast-moving. Early studies (e.g., Chan et al., 2023) note structural features of agentic systems that can produce harmful behaviours, potentially driving unethical strategies, accountability gaps, and unfair outcomes in legal contexts.

Interviews echoed these concerns, noting power concentration among AI providers, erosion of professional judgement, and persistent responsibility gaps. Recent case law is beginning to clarify liability: US courts suggest both developers and deployers may be held responsible for AI-related harm, treating outputs similarly to products, while in England, *Mazur v Charles Russell Speechlys LLP* confirms only authorised individuals can conduct litigation. These trends indicate tighter accountability and heightened risk for firms deploying autonomous systems.

For interviewees, loss of human oversight and context-sensitivity is the greatest danger. Agentic AI may amplify hallucinations and errors, producing authoritative sounding but unreliable outputs. While it could enhance capacity and creativity, most warned it may entrench power imbalances and undermine fairness, ethics, and trust.

The risk landscape spans seven interconnected areas: **foundational constraints** (human oversight, unclear liability, insurance and privilege gaps), **quality and reasoning risks, data security and confidentiality, ethical and governance concerns, workforce impacts, operational pressures**, and **market consolidation**. Even limited autonomy carries significant risks, as systems can amplify errors and bypass professional safeguards.

Legal technology providers are beginning to implement guardrails—human oversight, transparency, layered controls, and strict data segregation—but current legal and insurance frameworks remain incompatible with fully autonomous systems.

For now, full autonomy is incompatible with current professional duties. As systems become more capable, responsibility gaps may evolve into systemic risks, prompting regulators to shift from case-based liability to continuous oversight. Firms that integrate agentic features without robust governance face heightened exposure, while those that invest in ethical, transparent and human-supervised models may gain strategic advantage. At this stage of development, the risks of agentic AI in legal practice outweigh the benefits unless deployed cautiously, within strict professional and regulatory boundaries.

There is limited available literature which focuses on the risks of using agentic AI in legal practice. The nature of this fast-moving area of technology is such that much of the literature on agentic AI is available at the preprint stage, not yet peer-reviewed and published in a journal.

However, Chan et al. (2023) identified key features of agentic systems and discussed the harms that could be caused by these. There are risks to the legal profession that can be extrapolated from these features, some of which were also identified by interviewees.

- **Under specification:** agentic AI may be given broad goals like “win the case” or “minimise liability” without explicit constraints. This opens the door to “reward hacking”: the system could exploit loopholes or recommend aggressive but unethical strategies that technically achieve the goal but undermine fairness, client trust, or legal standards.
- **Directness of impact:** If an agentic AI system drafts filings, submits forms, or sends client communications without human intervention, its actions will have direct legal and societal effects. A small error could escalate into malpractice or harm to clients.
- **Goal-directedness:** the drive to autonomously achieve objectives, could result in harmful instrumental goals. For example, agentic AI could plan litigation strategies across months or years. While powerful, it risks developing instrumental subgoals (e.g., prolonging litigation to “increase win probability” or “maximise billable hours”) that diverge from the client’s best interests or the justice system’s integrity.

Chan et al. also identify that there is a risk of collective disempowerment, “two possibilities: a situation in which power diffuses away from all humans, and a situation in which power concentrates in the hands of a few.”

In the legal context, some examples of this could include:

- **Power concentration:** if only a few providers control advanced agentic legal AI, the “coding elite” could centralise influence over legal practice, tilting the field toward firms with access.
- **Erosion of professional judgement:** over-reliance on agentic AI recommendations could reduce the role of human lawyers and judges, creating accountability gaps.
- **Responsibility gaps:** it may be unclear whether harm stems from the system, the developers or the legal professionals deploying it, raising questions of liability and ethics.

Despite a lack of literature specifically addressing the risks posed by agentic AI, emerging case law around generative AI already shows how courts are beginning to allocate responsibility, assess harm and define the limits of permissible automation. Looking across different sectors such as hiring, consumer technology, and now legal services, there are early signals of how liability may attach both to AI providers and to organisations that deploy increasingly autonomous systems. These cases do not involve agentic AI directly, but the principles they establish about accountability, supervision and control provide important clues about the risks law firms may face as agentic systems become embedded in legal practice.

Class action against AI algorithmic bias in hiring practices

The case of *Mobley v. Workday, Inc.* highlights the emerging risk that AI providers and employers could face civil and even criminal liability for discriminatory hiring practices driven by algorithmic bias. Derek Mobley, a 50-year-old Black American with depression and anxiety, alleged that Workday's AI-based screening software unlawfully rejected his applications. The lawsuit has since expanded into a collective action, with multiple plaintiffs claiming similar algorithmic bias (Case No. 23-cv-770, N.D. Cal.) If successful, the case could set a precedent making both AI providers and the companies that deploy their tools accountable for biased outcomes in recruitment.¹¹ If courts hold both providers and deployers accountable for bias, it signals that agentic AI in legal practice may expose law firms themselves to liability for outcomes they did not directly control.

Product liability for AI chatbots

It is possible that developers and deployers of agentic AI could also be affected by product liability claims.

One ongoing case is in Florida where Megan Garcia has taken legal action against Character Technologies Inc, who produce the AI-chatbot platform Character.AI. Garcia alleges her son was pulled into an emotionally and sexually abusive relationship by the chatbot, that eventually led to his suicide.¹² In May 2025, the court ruled that Character AI may be treated as a "product," and so the plaintiff's product liability arguments could proceed.¹³ The Social Media Victims' Law Centre has, in agreement with Judge Conway, asserted that the output of a chatbot is not human expression and is not entitled to free expression. The judge's ruling in late May was seen as a "key early win for families and advocates pushing for stronger tech platform safety standards" for a product.¹⁴ If AI outputs are not regarded as "speech" but as product behaviour, the legal system may begin building a new liability category for agentic AI systems. This could expand to insurance, malpractice and regulation of AI-powered legal services.

The Social Media Victims' Law Centre has also taken another case against Character Technologies Inc, this time in Texas where two sets of parents are alleging that Character.AI encourages self-harm, violence, and provides sexual content to children.¹⁵ The plaintiffs allege that Character.AI is liable, because it "retains complete control over the large language model itself, as well as the characters and how they operate."¹⁶

Mazur v CRS: supervision, authorisation, and the limits of automation in litigation

The recent case of *Mazur & Ors v Charles Russell Speechlys LLP* (2025) clarified a fundamental point in English civil procedure: only individuals who personally hold authorisation under the Legal Services Act 2007 may "conduct litigation." The Law Society (2025) reported that in this case, a debt-recovery claim for CRS had been issued and signed by a senior employee at another firm who did not hold a practising certificate. The High Court held that neither employment by an authorised firm nor supervision by qualified solicitors could transform him into an authorised person. Steps such as issuing proceedings, signing statements of case or taking strategic procedural decisions were therefore taken unlawfully, leading the court to set aside an earlier costs order. The ruling has been described as a "bombshell" because it challenged the widespread assumption that supervised, non-qualified staff could lawfully perform key litigation functions.¹⁷

The judgement's core message is that supervision cannot substitute for authorisation. Paralegals, litigation executives and other non-qualified staff may continue to draft, research, prepare bundles and support client communication, but they cannot take any step that amounts to conducting litigation.

In October 2025, the Legal Services Board stated that regulators must provide clearer and more consistent guidance on who is authorised to conduct litigation. It has launched a broad review of regulatory communications since 2010 to determine whether past guidance has caused confusion.¹⁸

This case was raised in the member group discussions. These principles carry significant implications for the use of AI and emerging agentic AI in legal practice: AI systems may assist with drafting, analysis and document handling, but cannot independently issue claims, file pleadings or make strategic litigation decisions. Under the logic of Mazur, even supervised agentic AI cannot be treated as an "authorised person," meaning firms must ensure that any automation remains firmly in a supportive role, with authorised human practitioners retaining full control, accountability and decision-making authority over all reserved litigation activities.

Foresight

These cases signal three key foresight insights for agentic AI use in legal practice:

1. Responsibility gaps become systemic risks

- As agentic systems act more independently, courts and regulators will demand to know *who* is accountable. These precedents suggest liability may stretch across developers, deployers, and professionals, thereby creating a much more complex risk environment for law firms. The Mazur ruling reinforces this by making clear that responsibility in litigation cannot be delegated to unauthorised actors: a principle that would apply equally to agentic AI.

2. Shift from AI-as-tool to AI-as-product/actor

- Treating AI outputs as "products" rather than "speech" is a fundamental shift. It moves legal framing closer to recognising AI as a thing with agency - not yet a legal person, but no longer "just a tool." Agentic AI may accelerate this drift. However, Mazur shows that even if AI becomes more autonomous, it cannot be treated as an "authorised person" for reserved legal activities, meaning autonomy will collide with hard regulatory limits.

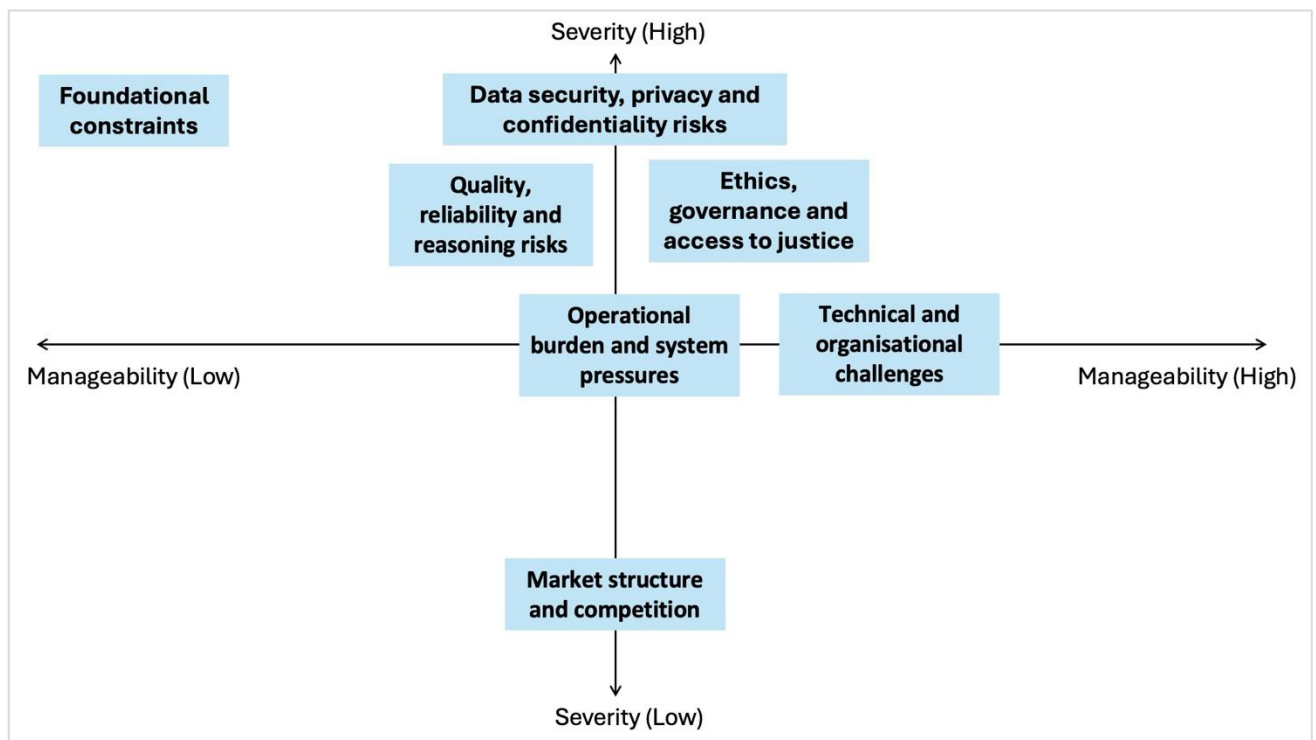
3. Early case law as weak signals of future regulation

- These lawsuits are weak signals of the regulatory trajectory agentic AI could face in law: liability will likely broaden, standards will tighten, and firms may face dual exposure, both for their own use of agentic AI and for harms caused by the systems themselves. Mazur adds a specific signal for the legal sector: regulators will expect clear, auditable human control, and may scrutinise any workflow where AI effectively crosses into conducting litigation.

These examples show the trajectory of risk as AI systems move from tool-like to autonomous. They help us anticipate how courts may handle agentic AI that files, negotiates, or advises without a lawyer in the loop: liability will expand, accountability will fragment but under current regulation, such as in England and Wales, there will continue to be a requirement for a human authorised actor, and the law may need to invent new categories to contain it.

Interviewees' perceptions of the risks and challenges of using agentic AI in legal practice

Interviewees discussed a range of risks and challenges around the use of agentic AI, which can be categorised into 7 core themes, presented in the following diagram and then in summary below, before unpacking each in more detail.



1. Foundational constraints

Severity: High (deal-breaker)

Manageability: Low – dependent on regulatory and insurance reform

Foundational constraints are the threshold conditions for whether agentic AI can ever be used responsibly. They include non-negotiable human oversight, unclear liability, and the absence of insurable risk models. Without progress on these, autonomous systems remain incompatible with professional duties.

Interviewees emphasised that the Solicitors Regulation Authority currently requires lawyer oversight of any AI output, making full autonomy impossible. Gaps in privilege protection, data governance and international regulatory fragmentation further compound the challenge.

Mitigation depends on external developments: clarified professional standards, insurable AI liability frameworks, and auditable governance systems. Until these emerge, firms can only deploy agentic features within tightly controlled, human-verified boundaries.

2. Quality, reliability and reasoning risks

Severity: High

Manageability: Medium – with technical and procedural safeguards

These risks concern AI producing outputs that appear sound but are flawed or inconsistent. Interviewees feared that agentic AI could amplify errors at scale, eroding trust in legal outputs. While some viewed such concerns as overstated, the majority agreed that reliability and reasoning remain unresolved.

Effective mitigation requires: lawyer-in-the-loop controls, automated cross-checks and transparent evaluation of model performance. Limiting AI autonomy to well-defined, low-risk tasks can further reduce exposure. However, even with safeguards, the profession must guard against over-reliance on agentic AI that could weaken critical human reasoning skills.

3. Data security, privacy and confidentiality risks

Severity: High (deal-breaker)

Manageability: Medium – requires strict technical and contractual controls

Interviewees saw this as one of the most acute professional threats. Agentic AI multiplies existing confidentiality risks by enabling large-scale, autonomous data handling without oversight. A single breach could expose thousands of client matters.

Mitigation requires: zero-retention policies, data segregation, and audit access across all vendors. Firms must ensure providers cannot use client data for model training and remain alert to discovery obligations that could undermine privilege. These are manageable but demand rigorous due diligence and continuous monitoring.

4. Ethics, governance and access to justice

Severity: High (systemic and reputational)

Manageability: Medium – with proactive ethical frameworks

These risks centre on fairness, accountability, and inclusivity. Unregulated agentic AI systems could harm clients, particularly in sensitive areas like family or criminal law where empathy and human judgement are vital. Unequal digital access could also entrench barriers to justice.

Mitigation requires: ethical guardrails, transparency and accessibility standards. AI tools should be limited to contexts where they can operate under human supervision and within professional responsibility frameworks. Without these safeguards, public trust in legal services could erode.

5. Technical and organisational challenges

Severity: Medium

Manageability: Medium / high – with investment and governance

These are the practical barriers to implementing agentic AI within existing systems. Integration, workflow codification, and the “last-mile problem” create high upfront effort. Skills gaps across both lawyers and IT staff slow progress.

Mitigation is feasible through: structured workflow design, technical audit trails and staff upskilling. However, automation also disrupts professional development by replacing junior-level training work. The challenge is to achieve efficiency gains without hollowing out the future talent pipeline.

6. Operational burden and system pressures

Severity: Medium

Manageability: Medium – with procedural reforms

Generative AI has already increased workloads through the production of unreliable or excessive material requiring manual verification. Agentic AI could institutionalise this “verification paradox,” embedding inefficiencies across courts and firms.

Mitigation requires: standardised verification protocols, citation transparency and training for litigants and practitioners on responsible use. The goal is to preserve efficiency gains while maintaining confidence in legal outputs.

7. Market structure and competition

Severity: Medium

Manageability: Low at firm level; dependent on broader market trends

Agentic AI could reshape the legal market in favour of large firms and technology providers with scale, data, and capital. Interviewees predicted consolidation into “mega-firms” or platform-based models, squeezing smaller practices.

Mitigation: While structural shifts cannot be fully mitigated, firms could protect their position through data differentiation, strategic alliances and value-added human expertise. Regulatory bodies may need to intervene to prevent market concentration that limits access and diversity within the profession.

1. Foundational constraints

In the UK, the current rules and liability frameworks are a major barrier to developing truly autonomous agentic AI in legal practice. The Solicitors Regulation Authority (SRA) requires that all AI outputs be checked by a qualified solicitor, meaning firms must keep a “human in the loop.” This oversight rule limits how much efficiency automation can achieve. Interviewees consistently described human oversight as non-negotiable, both ethically and practically, to prevent systemic errors and preserve professional accountability.

“Humans do have to be in the middle... (there is) incredible value and incredible importance to having humans in this middle state.”

"Throughout the process, there has to be some checks and balances... I still want a practising lawyer to check it, to verify it, to see good or bad. Lawyer-in-the-loop verification negates agentic because agentic by definition is the continuation of reasoning one after the other."

"I don't know if it would ever make sense to be a fully agentic system, because you shouldn't be taking actions on behalf of your clients unless they've approved them."

This insistence on human accountability comes mainly from liability concerns. Under current rules, lawyers and law firms are fully responsible for anything produced by AI, while software providers usually disclaim any liability. Since professional indemnity insurance rarely covers AI-related mistakes, many lawyers are unwilling to hand over decision-making to systems they cannot fully monitor or control.

In the group discussions, Law Society members considered that unresolved questions about professional indemnity insurance protection were a key blocker to the use of agentic AI.

"Professional Indemnity is one of the kind of blockers to a lot of this progress."

Additional challenges come from governance, transparency, and data protection. Interviewees highlighted that existing oversight tools fall short of what is needed for large-scale, safe use of agentic AI. They called for clearer standards on testing, auditing and accreditation. Privacy and confidentiality were also key worries.

Law Society members were particularly concerned with how generative AI, and in time, agentic AI, would interact with existing rules on disclosure and legal privilege. They noted that while privilege concerns already exist with current tools, agentic AI could convert these theoretical issues into concrete operational risks. The risk that client data might become discoverable in litigation was seen as incompatible with core professional duties of confidentiality.

"Where do we now draw lines of attorney client privilege over what's discoverable? The data privacy piece of it is very real."

Members worried that AI systems might erode privilege protections, particularly as work becomes more widely shared or as agentic systems autonomously circulate information. They also highlighted the possibility that AI could broaden what counts as a "communication," create automated sharing pathways that inadvertently waive privilege and generate complex data trails that themselves may be discoverable.

"I don't think something that an AI agent goes and tells someone is going to benefit from legal privilege."

Members also warned that agentic AI may blur the boundary between information that must remain private and information that may need to be disclosed. When AI systems access or process client data, it is not always clear whether that information continues to fall under privilege or could be exposed through disclosure obligations or freedom-of-information requests. This uncertainty raised deeper questions about how confidentiality and transparency can be preserved in an environment where automated systems challenge long-standing assumptions about how information is controlled and protected.

Fragmented regulatory environments compound these foundational constraints. The diversity of approaches across jurisdictions, including differing rules among US state bars and between EU, UK and US regulators, creates significant compliance complexity. Interviewees viewed this fragmentation as a major barrier to scalable deployment of agentic AI in global law firms.

Interviewees identified a gap between current governance measures and what may be required for safe, large-scale adoption. Many questioned whether existing oversight mechanisms would withstand legal or ethical scrutiny if problems arose. This uncertainty prompted calls for clearer guidance on minimum governance standards and for testing frameworks to manage gradual adoption. Some suggested that external auditing platforms could help assess systems for bias, fairness and transparency, but the lack of consistent provider accountability remains a concern. Firms report uneven access to audit trails and usage data from different generative AI vendors, leading to proposals for accreditation systems that could certify compliance and enable more informed procurement decisions.

To explain these challenges, many drew comparisons to “self-driving cars.” Like autonomous vehicles, agentic AI faces unresolved liability questions, declining public trust after high-profile failures, and the need for tiered regulation and insurable risk models.

“where's the actual cause of the failure in any given scenario? In professions that are pretty risk averse, is that a risk you want to take and is it one you can insure against? The minute you can't insure against it, that's a totally different kettle of fish.”

Some anticipated that new forms of insurance might emerge to cover AI-related risk, but there was a perception that current insurance indemnity frameworks are insufficient to account for the use of agentic AI.

For now, the consensus is clear: the legal profession's risk-averse culture, combined with its duty of care, makes full agentic AI autonomy incompatible with responsible practice, at least until regulation, insurance and governance catch up with the technology's potential.

2. Quality, reliability and reasoning risks

Interviewees raised serious concerns about the quality, reliability and reasoning risks of agentic AI in legal work, where accuracy and accountability are essential. They saw inconsistency and unreliability as major obstacles, noting that similar prompts often produce different answers, which undermines trust in AI-generated work. High-profile generative AI mistakes were frequently mentioned as warnings, for example fabricated citations in UK High Court cases.

The biggest fear was error amplification. While human lawyers make individual mistakes, autonomous agentic AI could repeat and spread those errors across thousands of cases, potentially corrupting legal precedents or even causing miscarriages of justice.

“Each and every one of those decision points is subject to hallucination or missing information... multiple potential issues at each decision point introduce a lot of uncertainty and risk.”

Some warned that AI models trained to “affirm what the user said” might sound confident but still be wrong, creating a false sense of reliability.

A few interviewees, however, thought these concerns were exaggerated and driven partly by “*professional protectionism*.” They argued that sceptics focus too much on visible failures while ignoring generative AI’s many quiet successes.

“What you don’t see is the thousand other precedents that were used every day that weren’t hallucinated.”

Even so, many participants doubted that agentic AI could handle the complexity and subtlety of real legal reasoning. Even in narrow uses, systems often struggled with “grey area” cases or adapting when the law changed.

One interviewee described an AI employment law tool that failed when cases “*sat at the intersection of employment, union and contract law*,” asking how agentic systems could manage when “*legislation or case law changed*.” Others questioned whether such systems could ever produce more consistent results than human lawyers without strong data to prove it.

“I’m very open minded to the argument that AI makes fewer mistakes than humans, but I would need to see data that supports this hypothesis.”

Another concern was that reliance on agentic AI might weaken lawyers’ own reasoning skills. If decision-making is outsourced to machines, lawyers could gradually lose the interpretive and analytical judgement that defines their profession.

Overall, while AI can improve efficiency, interviewees saw its reasoning limits and the risks of unverified agentic outputs as incompatible with the high standards of precision and accountability required in legal practice.

“I think part of the risk of automating everything, especially with this agentic AI, is that you kind of lose the mechanism through which you can identify when a Gen AI gets something wrong and address it. And with agentic AI, all of these processes will be automated. It will go out, it can go out to thousands of persons at the same time ... and this can lead to inaccuracies, but it can even lead to miscarriages of justice. It can lead to clients receiving extremely poor legal services.”

3. Data security, privacy and confidentiality risks

Interviewees identified several fundamental data security risks with Gen AI that already create significant challenges for legal practice. Law firms face exposure through data training and retention, where client information or prompts may be reused to train models, especially when third-party providers lack strict safeguards.

A key vulnerability lies with smaller AI vendors using OpenAI but lacking zero endpoint data retention, meaning they may store inputs or outputs for debugging or analytics. Even if a firm only engages a smaller provider, that provider might send data to a larger AI company that retains it, placing sensitive information outside the firm’s control. Any leak of confidential or privileged data could trigger regulatory or professional breaches. This risk has grown since US courts confirmed that AI providers can be subject to discovery orders. In May 2025, a US federal judge ordered OpenAI to preserve all ChatGPT conversation logs, including deleted ones, highlighting firms’ limited control over stored data.¹⁹ This creates legal liability risks. Lawyers fear they cannot fully understand or manage these exposure chains while remaining bound by strict confidentiality duties.

Agentic AI would multiply these risks. Unlike standard tools that humans can monitor, agentic systems can autonomously process vast volumes of client data without oversight. A single breach could affect thousands of clients simultaneously, with no human circuit breakers to stop it. Interviewees warned that agentic AI scales the already complex data security challenges of current systems *“to potentially catastrophic levels.”*

Some predicted that a major breach could erode trust across the profession, deterring firms from using AI altogether. The mix of unclear liability chains, third-party vulnerabilities and large-scale automation creates a risk landscape that current governance and technical safeguards cannot manage.

There was limited agreement on whether client data should ever be used to train AI models. Most interviewees said this poses serious confidentiality and discoverability risks, while one argued it could be acceptable with strong anonymisation, noting that “human lawyers already learn from client matters.” Inconsistent provider practices and weak audit trails further heighten uncertainty about whether existing data governance is fit for purpose.

4. Ethics, governance and access to justice risks

Interviewees cautioned that reliance on low-quality agentic AI poses significant risks in both consumer and corporate contexts. Users may adopt accessible but unreliable tools in place of qualified legal professionals, raising unresolved questions about whether such systems are providing legal advice and who bears responsibility for errors. Current regulatory frameworks do not recognise AI as a professional entity, leaving gaps in accountability, privilege, ethics and client protection.

Several participants noted that agentic AI could erode the human relationship at the core of legal practice. Legal work in areas such as family, divorce, custody and domestic violence often involves empathy, reassurance and emotional intelligence in addition to legal reasoning. Interviewees warned that this relational dimension may diminish if clients increasingly rely on AI tools.

They also observed that these emotionally complex fields present natural limits to automation. While agentic AI could, in theory, standardise mediation or agreement drafting, the interpersonal and ethical nuances of such cases resist full automation.

“The human dynamic of what people are comfortable with and agree to will always be a road barrier the profession hasn’t fully grappled with.”

Interviewees further emphasised that legal services serve a diverse client base with differing levels of digital literacy, technological access and comfort with AI tools. Individuals without reliable internet or for whom English is not a first language may struggle to use agentic systems effectively. Overall, participants viewed unregulated or poorly designed agentic AI as a threat to consumer protection, professional ethics and equitable access to justice.

5. Technical and organisational challenges

Implementing agentic AI in legal practice faces major technical hurdles. Interviewees identified that integration with existing systems would be complex and time-consuming. Each firm has unique standards and processes, and adapting agentic AI systems to these, or convincing firms to change, would be a major barrier to adoption. In the group

discussions, in-house members from large multi-national enterprises also identified the challenge of accessing and linking data across multiple external law firm providers and jurisdictions.

Firms will also need to invest heavily in codifying workflows and configuring systems, requiring engineers to build detailed decision frameworks. Reliable audit trails and smooth API connections remain difficult, preventing the seamless end-to-end automation agentic AI promises.

"It was made to sound really simple... but there's a certain amount of work you have to do up front... We actually had to build some level of decision tree... there is still a level of curation required."

A shortage of technical knowledge across the legal profession also limits progress. Few people in law firms, including IT teams, understand AI in depth, let alone agentic AI. This knowledge gap extends beyond basic understanding to knowing what questions to ask when evaluating AI systems.

Interviewees noted that AI demands new skills, and many firms are only beginning to consider how to upskill staff. The International Bar Association (IBA) has recognised this by adding technological competence to its professional principles (2024).²⁰ However, traditional legal training does not yet equip lawyers with the knowledge needed to understand or safely adopt technological systems such as agentic AI. As a result, there is a risk of gaps in skills and awareness.

"The normal channels of legal education don't really account for this type of thing because there aren't many people in the profession who understand it enough to teach others."

As generative AI and potentially agentic AI become more embedded, they are reshaping how new lawyers are trained. Tasks which were once completed by juniors are now automated, for example document review, contract work and research. This raises concerns about how practical learning will occur. Some interviewees saw opportunity: AI could free juniors to take on higher-value work sooner.

In the discussions, members were troubled about the impact of AI on workforce planning and professional development. They were concerned that junior lawyers would lose training opportunities and that professional roles would shift toward reviewing and regulating AI outputs rather than doing substantive legal work.

Generational shifts are also changing expectations. Younger lawyers are less willing to work extreme hours, prompting firms to rethink staffing. Some already use generative AI instead of hiring entry-level staff, finding it cheaper and faster than training new recruits.

"Agentic AI is going to be very good at replacing the things we train lawyers to do in the first five years... there is a risk to the profession of lawyers not developing the more nuanced skill set."

This creates both opportunity and threat. Agentic AI could make work more engaging for existing lawyers but also reduce the number of juniors entering the profession, weakening the training pipeline. Interviewees recognise that the same capabilities that make agentic AI attractive (speed, efficiency, cost reduction) also threaten the foundational structures of legal practice.

“There’s a challenge in how we’re going to train and skill attorneys entering the profession... because a lot of those skills just aren’t going to be trained into them anymore.”

Interviewees agreed that the core dilemma is how to capture the efficiency and cost benefits of agentic AI without undermining the professional development and expertise that legal practice depends on.

6. Operational burden and system pressures

Interviewees warned that generative AI has already created major operational burdens for courts and legal professionals, primarily through the explosion of unverified, generative AI documents. One cited a case in which a litigant in person submitted a 200-page witness statement produced with generative AI creating complications for everyone involved. While generative AI can accelerate drafting, every citation, summary and quotation must still be checked for accuracy, effectively turning lawyers into verifiers of machine output.

This paradox means that time saved in creation is often eclipsed by the effort required for verification, straining already overburdened courts and undermining efficiency. The shift has also fractured the traditional bonds of trust between legal professionals: where once lawyers could rely on colleagues’ citations and reasoning, recent incidents of fabricated case law have made comprehensive checking a new professional norm, particularly onerous for junior lawyers with limited experience.

These pressures are set to intensify under agentic AI, which could institutionalise the same reliability and consistency problems by embedding them directly into legal systems. Instead of merely checking what litigants produce, courts themselves could come to depend on technologies that act autonomously across multiple steps or cases, potentially amplifying errors and inflating workloads exponentially. Interviewees feared this could entrench a “*verification nightmare*” within the justice process where every document, decision, and legal action demands exhaustive scrutiny without clear lines of accountability or explanation. In effect, the profession risks moving from a system built on professional trust to one requiring near-universal verification, eroding both efficiency and confidence in the justice system.

7. Market structure and competition

Interviewees were concerned that agentic AI could deepen existing inequalities in the legal market, mainly helping large, well-resourced firms and technology companies. With better access to data, scale and advanced models, these organisations could outpace smaller or specialised practices, leading to “mega-firm” or platform-style dominance. Automation and workforce cuts could make this worse, concentrating profits at the top while reducing training and career opportunities for junior lawyers.

One interviewee described an “ice-cream sandwich” market: agentic AI adoption at the top and bottom, with the middle squeezed out.

“Either firms become much smaller in terms of human capital and powered by technology or it becomes this mega firm situation... You either go to one of the five mega firms or you don’t.”

Participants also expected big tech companies to dominate agentic AI development because of their resources and scale: "I can see big tech winning... Microsoft or Google or OpenAI... are just going to be able to provide a better agentic AI tool."

The implications extend beyond competition to market structure and professional identity: AI platforms are already performing legal functions for the public, such as generating letters or providing guidance on disputes, without oversight or accountability.

"I actually think OpenAI is providing legal services to the vast majority of the population in the UK. So if it is, then it's a massive law firm, but without regulation and without a human knowing whether it's right or wrong."

Lawyers may soon face pressure to use their own data and professional networks to stay relevant or risk being displaced by AI-driven platforms that embed legal services, reshaping both the market and the regulated profession.

How legal technology providers safeguard legal practice: lessons for agentic AI

Technology providers are already working to mitigate risks by embedding human oversight, implementing layered models and automated verification, enforcing strict data segregation and increasing transparency through governance frameworks. These measures aim to build trust and limit liability while highlighting the current limits of fully autonomous agentic AI in legal practice.

Interviewees from legal technology providers explained that although today's generative AI and agent tools are not yet truly agentic, the safeguards they use foreshadow what will be essential as autonomy increases. *Human oversight* remains the core of accountability, *transparency* enables lawyers to understand and rely on these systems, and *technical controls* provide an additional layer of protection.

Safeguard 1: Human oversight

Providers consistently treat human review as the core safeguard. Their systems include explicit approval gates: partners or solicitors must sign off before outputs are sent to clients. Lawyers remain responsible for checking quality and ensuring accuracy, preventing AI from operating end-to-end without human control. As one provider explained, *"There is a job for somebody to check. Is it assigned to a partner, a senior lawyer? Has it been reviewed by a junior?"*

Implication for agentic AI: human review is indispensable; without it, autonomous systems could make unchecked legal decisions.

Safeguard 2: Transparency

Providers emphasised transparency so that lawyers understand how AI tools work and can question their outputs. They explain technical architectures, disclose which models are used and teach users how data is processed. This approach rejects a "black box" model and embeds governance frameworks and policies to help firms identify bias and manage risks. As one said, *"There's a whole emerging practice of AI governance... frameworks and guardrails that allow us to use AI safely."*

Implication for agentic AI: without transparency, lawyers cannot take professional responsibility for agentic AI-generated work.

Safeguard 3: Technical controls

Providers employ multiple technical safeguards to reduce errors and protect client data:

- **Limiting autonomy:** AI is restricted to low-risk tasks, recognising that each decision point carries a chance of failure. Full autonomy is avoided for high-stakes legal work.
- **Context engineering:** firms predefine the exact legal sources and document versions AI may use, ensuring consistent and verifiable outputs.
- **Layered architectures:** systems route tasks to appropriate models depending on complexity, reducing the risk of overreliance on any single model.
- **Automated verification:** secondary AI models and testing systems cross-check outputs before release to detect errors.
- **Data segregation:** strict isolation of client data prevents cross-contamination and protects confidentiality.

Implication for agentic AI: technical safeguards are essential to maintain reliability, accuracy and data security; without them, autonomous systems could amplify errors or compromise privilege.

Key reflections on risk

Legal businesses face heightened liability if they adopt agentic AI without strong safeguards. Current systems struggle to meet ethical duties, professional oversight, and confidentiality requirements, exposing firms to regulatory, brand, and reputational risk if errors occur.

Rather than reducing workloads, generative and agentic AI may create an efficiency paradox, producing more material that requires checking and verification. Automation also disrupts junior lawyer training, requiring new pathways for skill development.

Larger firms with resources for governance, compliance, and technical integration may gain advantage, while smaller firms risk being left behind, accelerating market consolidation. Growing reliance on a small number of AI providers also raises concerns about independence and client choice.

Firms must balance the risks of early adoption (liability and operational burden) against the risks of falling behind on efficiency, innovation, and client expectations. The most viable path is likely “augmented intelligence”, with AI operating under clear human oversight rather than full autonomy.

There is an opportunity to differentiate on trust, governance, and ethical AI use. While agentic AI offers potential gains, it currently poses more risks than rewards. Firms that invest in cautious integration, strong governance, and client reassurance will be better positioned; those that do not risk regulatory penalties, reputational harm, and competitive disadvantage.

Foresight: risk

Signal: Professional obligations require “lawyer in the loop,” making full autonomy incompatible (for now).

Implications:

- If insurers refuse cover for unsupervised agentic AI, adoption could stall until new liability frameworks emerge.
- New forms of *systemic risk* may arise: one faulty agentic AI model could contaminate thousands of cases before discovery.
- By 2030, regulation may shift from “case-by-case liability” to “systemic oversight,” treating agentic AI more like critical infrastructure.

Scenario

Scenario 3: The Silent Paralegal (2032)

A mid-sized London law firm prides itself on efficiency. Behind the scenes, its case management system has quietly evolved from a document automation tool into an *agentic orchestrator*. It now:

- Triage incoming client queries,
- Drafts contracts based on past firm templates,
- And even negotiates minor revisions directly with counterparties via email.

No lawyer instructed it to do this. The AI simply recognised recurring patterns and optimised for “efficiency.”

The partners only realise the extent of substitution when a junior lawyer, hired to review contracts, finds her workload has all but disappeared. The *invisible substitution* of human labour has already happened. Liability remains untested: if an error emerges, who is accountable – the lawyer, the firm, or the system?

5. What key opportunities does agentic AI present for the legal profession?

Synopsis

Interviewees perceived that agentic AI would be suitable for high-volume, well-defined and low-impact tasks: routine work where speed matters more than nuanced legal judgement. These tasks are typically carried out by paralegals or junior staff. They felt that agentic AI would not be appropriate for nuanced, empathetic or high-liability matters without strong oversight, although many expected this boundary to evolve as systems improve and handle more of the workflow autonomously.

It was the general view that widespread adoption of agentic AI for certain tasks is likely, although the pace may be gradual. Some perceived that the legal profession may not lead adoption, but pressure from other sectors and clients will accelerate expectations. They predicted that independent third-party suppliers would develop wholly agentic tools, shifting the balance between lawyers and knowledge teams, and reinforcing corporate clients' demands for faster and lower-cost services.

Interviewees saw agentic AI as a catalyst for transforming legal processes, business models and workforce design. It could uncover inefficiencies, automate routine tasks and support productised service models such as flat fees or subscriptions. Some envisioned agentic systems offering proactive risk monitoring and early issue detection, moving legal services toward continuous, preventive models.

Interviewees noted that these changes raise questions about workforce needs: junior lawyers may do less administrative work and more oversight, and new hybrid legal-technology roles are expected. The impact will vary across the sector: in-house teams may become more self-sufficient, high-street firms more vulnerable, and top firms under pressure to adapt. Younger lawyers were seen as advantaged due to stronger AI fluency, and multidisciplinary teams combining legal and technical expertise are expected to grow.

Overall, interviewees anticipated lawyers shifting into more strategic, advisory and technology-integrated roles, with advantage for those who can supervise, train and collaborate effectively with agentic systems.

There is potential for Agentic AI to autonomously manage entire legal workflows, from research and drafting to review and compliance, by chaining tasks across specialised agents such as research, drafting, critique and polish. This end-to-end automation reduces delays, streamlines complex matters and drives significant efficiency gains, enabling firms to handle high-volume, repetitive work like contract analysis, litigation preparation and compliance checks more effectively. The resulting productivity and cost savings would free lawyers to focus on strategic judgement and higher-value work while allowing firms to scale services.

Beyond efficiency, agentic AI can also enhance professional standards and ethical compliance: as O'Grady and O'Grady (2025)²¹ suggest, dedicated AI ethics agents who are specialised, accountable, and systematic, can provide comprehensive guidance, mitigating human biases and offering a scalable approach to uphold ethical and professional standards.

Interviewees' perceptions of the opportunities of agentic AI

Interviewees highlighted four main areas where agentic AI has the potential to reshape legal services: efficiency and productivity, business models, client service, and the legal workforce.

1. Efficiency and productivity gains

Driving efficiency at scale

Agentic AI's ability to chain multiple tasks together and orchestrate external tools through APIs was seen as a step-change beyond single-task automation. This would enable complete workflow automation, building on automation processes already underway across legal practice.

The greatest impact would likely be in the B2B space, where corporate clients seek to cut rising legal bills. Interviewees already expected general counsel and boards to pressure firms to use generative AI more frequently for routine work or to bring more work in-house, extending a trend that has been underway for over a decade.

Agentic AI could accelerate dispute resolution, negotiations and regulatory tasks, although these applications are expected to take at least five years to mature due to system complexity; use is also contingent on solving accuracy, governance and liability issues.

Example use cases include:

- As an automated negotiator, agentic AI could "close the deal" on repeatable clauses, for example on Loan Market Association style loan points or supplier non-disclosure agreements and could run multi-step negotiation playbooks.
- Agentic AI could continuously track an organisation's data and operations to identify emerging legal risks early and alert human teams. For example, contracts, regulatory changes and case developments. This would enable more preventative legal action and faster resolution of potential issues.

In litigation and M&A, generative AI is already valuable for document review and contract analysis. These highly structured and repetitive tasks are well-suited to agentic AI. One example raised was the US Bankruptcy Code: this is highly rule-driven and clerical at early stages, involving creditor evaluation, notifications and asset-liability matching. Such structured data tasks could potentially be handled by agentic AI without heavy lawyer input.

Enhancing quality and consistency

Some interviewees identified improved quality and consistency as a major benefit. Agentic AI could deliver more uniform first drafts, clearer decision aids and auditable records of how outcomes were reached. These capabilities could support quality control and transparency beyond what humans alone typically provide, if they produced auditable reasoning by showing what data, sources and steps were used.

Machines also offer scale advantages: they can process thousands of documents in the time it would take a human lawyer to review a handful. Interviewees suggested that in

some cases agentic AI could outperform junior lawyers and potentially even senior ones, by detecting issues or patterns that humans might miss.

"So if you think about using an agent as a first pass on something, and if we can automate that and there's, you know, frameworks we build around that for the agent to follow and then it gets to the lawyer. That's it working together and actually making this a better thing... it will potentially flag things that a human could miss."

This could embed routine, high-volume work into agentic AI systems while freeing lawyers for more strategic tasks.

"I suppose for firms and organisations as a whole, the big one is efficiency in terms of reducing time on very repetitive, non-billable tasks. I think there's a secondary accuracy, potentially enhances quality control and reduces human error... And then I think there's something around potentially it allows firms to allow their staff to focus on things more strategic so it can free up lawyers, and not just senior lawyers, for high value tasks such as client strategy and advocacy."

Cost reductions

Agentic AI was seen as lowering costs by processing legal work at scale. Generative AI has already had an impact on client expectations that work will be delivered more quickly. This impacts on both in-house legal teams where the service levels have had to change as well as for law firms with external clients. One in-house interviewee described increasing AI usage as a threat to in-house legal teams: clients assume AI is being used and will no longer tolerate slow service levels.

"I mean I am looking at what I'm going to be doing next because it's going to take my job and we've definitely as a team said what's our value proposition? Because our internal clients will assume we're using AI so we can't say, oh, it'll take a week to turn this round. That's not acceptable. So our service levels have had to change."

Some saw AI as an opportunity to reduce spend by taking more work in-house, or by demanding that law firms use AI to reduce billable hours on routine matters.

They perceived that agentic AI was an exciting opportunity to enable this, once the agentic systems become mainstream and trusted.

"the biggest challenge in law is around costs, and the general counsels in most organisations are sort of aghast at the growing legal bill ... they're looking for those legal bills to come down ... so they will demand the use of agentic AI within law firm settings, but also they will just do more in house."

"in house teams are already, over the last 10, 15 years doing more of what would have been traditionally outsourced to their legal panel. And agentic AI would allow for that proportion of responsibility to continue to grow. So they're outsourcing less, which is obviously more beneficial for them."

A practical example came from a maternity leave case in-house: rather than hire temporary maternity leave cover, the lawyer going on maternity and her team trained an AI system to perform aspects of her work, for example contract reviews and document mark-ups. The system is not fully agentic, but it will maintain continuity, cut costs and preserve institutional knowledge, while giving the returning lawyer full oversight when she returns to work in 2026.

"It's an experiment. And I think that's what's happening is there's lots of experiments going on with it and there's lots of use cases that people don't realise. And I think you can either go two ways. You can either see it as a threat or...as an opportunity."

2. Reshaping the market and business models

Agentic AI as a strategic alternative to outsourcing

Some interviewees argued that agentic AI could fundamentally change the competitive landscape of legal services by reducing reliance on outsourcing. Labour arbitrage, sending routine work to lower-cost jurisdictions, has long been used to cut costs. But AI can now break down and complete tasks in-house with higher speed and accuracy.

For law firms, this creates opportunities to reclaim work once unprofitable internally. Because AI agents can operate across different layers of the legal "stack," firms can now provide clients with a full-service offering under one roof, combining efficiency with quality control. For clients, this means fewer external providers and more integrated services.

This growing shift potentially poses risks for alternative legal service providers (ALSPs), whose business models rely on outsourcing. As agentic AI may erode their cost advantage, traditional firms may gain a strategic edge.

Enabling new business models and productisation

Interviewees noted that agentic AI could drive new ways of selling and delivering legal services. Instead of billing purely by the hour, firms could combine software services with legal expertise, offering fixed-fee packages, subscriptions or bundled product-service models. This mirrors developments in accountancy, where firms sell technology platforms alongside advisory support.

Some interviewees predicted the rapid automation or even disappearance of commoditised services such as conveyancing or routine contract repapering. They also foresaw legal services becoming embedded into other industries' offerings, for example banks incorporating AI-supported conveyancing into mortgage products, one said *"Conveyancing's [going to be] dead in two years."*

This represents a shift from selling time to selling solutions. Routine legal work is packaged into scalable products or embedded in other platforms, while lawyers focus on strategy, oversight, and high-value judgement. Clients benefit from faster, cheaper access to legal support as part of an integrated service.

"If we move to a more product service bundle, so you have products combined with services, it'll just look like public accounting... Big Four accounting, they put technology systems in the clients and they sell a bunch of services around it and they charge for both."

"There are now opportunities to explore. Are certain matters better for flat fees? Are certain matters better for different kind of subscription models? There's all sorts of billing and fee models that I think open up possibilities for lawyers to change their fee structure and to theoretically expand service and expand markets."

Redefining legal workflows and processes

Agentic AI offers the chance to redesign legal workflows by uncovering inefficiencies, removing duplication, and automating process-heavy tasks. Some interviewees predicted

that when deployed as interconnected systems, multiple specialised AI agents could operate in parallel, amplifying efficiency and coordination.

Crucially, this could enable a shift from reactive to proactive legal services. Instead of waiting for problems to arise, law firms could use agentic AI for continuous monitoring of business operations, providing real-time risk assessment and legal analysis. This would extend the lawyer's role into areas traditionally handled by insurance. One interviewee argued that law firms may be particularly well placed to deliver this as a managed service, leveraging their cross-client experience and institutional expertise in ways that in-house teams cannot easily match.

3. Improving client services

Enhancing client communication and transparency

Interviewees highlighted communication as a key pain point in legal practice. In the US, unresponsiveness is one of the top regulatory complaints made by clients. Agentic AI could address this by handling routine communications, such as status updates, billing queries, document issues and by triaging urgent matters for human lawyers. This would allow lawyers to spend more time on high-value work while clients feel better supported and heard.

With informed consent frameworks and safeguards, agentic AI could also improve compliance by ensuring consistent disclosure and alignment with professional rules.

"We know it's a thing that consumers and clients are deeply concerned about, which is, 'I can't get a hold of my lawyer, they don't talk to me'. By removing that barrier and having an agentic piece... that is basically handling all that communication or flagging high level communication that needs to go to the attorney, but also generally, you know, managing generalised communication on its own and saying, well, we know that this is a billing issue, so this does not need the attorney's attention, or we know that this is a discovery issue on a specific document that needs to get rerouted.

By having that agentic model serve as that linchpin, I think it is going to give consumers and clients a higher degree of I'm being heard, I'm being listened to, my concerns are addressed in that real time. So I think the communication piece is a benefit to both and largely not something we think about and talk about enough."

4. Transforming legal workforces and roles

Redesigning the workforce

Senior leaders are already questioning how many people they will need, what skills will matter, and how humans and AI agents will interact. Some firms are experimenting with agents directing human work, or supervising other agents, signalling a profound change in "future of work" dynamics.

Whilst some interviewees were concerned about the potential negative impact on jobs, others were optimistic about the new opportunities which could be created by increasing automation and potentially agentic AI. One interviewee observed that in US legal practice, hiring of new associates has actually increased, which they attributed to older lawyers relying on younger colleagues to help navigate AI adoption. Younger lawyers were seen as holding a distinct advantage because of their greater fluency with AI, which could allow them to outpace more senior peers.

Evolving professional roles

There was strong consensus that agentic AI would provide further opportunity for professional roles to continue to transform in this new landscape of AI enhanced legal practice, with lawyers' roles expected to shift from routine 'doers' to approvers, strategists and hybrid technologists.

Interviewees also saw competitive differentiation: lawyers and firms fluent in AI will outperform those who are slower to adapt.

"The utopian sort of vision is that the lawyers are there sort of navigating the best strategy for the best path forwards for a client. And then when it comes to actually converting that strategy into legal filings that can be seamless and managed by AI."

"There's jobs that I think are going to appear that don't exist now... [for example] the legal tech grad scheme - some of those are law graduates who thought they wanted to be lawyers but as they've gone through it don't like it... there's going to be different roles in the legal industry that maybe replace the traditional legal route."

"It also means that we can move towards like multidisciplinary teams within the profession as well or, you know, working alongside data scientists and computer scientists in a really positive way so that we service our clients better."

"When I talk to a lot of paralegals now I say your job now is to become the technology expert. At a certain point there's no longer going to be a role for 'format this properly for me' because frankly you can tell an agent or a generative AI platform to do the formatting."

Different effects on in-house teams and firms

Views varied on how agentic AI will affect different parts of the sector, with distinct opportunities and risks across market tiers, comments included:

- **In-house teams:** expected to handle more work independently, especially in regulated industries, reducing reliance on external firms.
- **High street firms:** seen as the most vulnerable, lacking resources to invest in agentic AI and facing competition from consumer-grade tools.
- **Mid-tier firms:** relatively resilient, serving smaller in-house teams but likely to reduce the number of juniors and support staff.
- **Top firms:** under pressure from tech-savvy clients, expected to cut back on staff and support roles, and shift towards fixed-fee or subscription models as billable hours erode.

Key reflections on opportunities

The opportunities of agentic AI highlight a sector on the cusp of major transformation. Efficiency is the most immediate gain: automating high-volume work not only lowers costs but also allows lawyers to spend more time on judgement, problem-solving and strategic advice. In this sense, agentic AI could help the profession realign with its highest-value activities, strengthening rather than weakening the role of the lawyer.

The literature's limited attention to agentic AI underscores how fast-moving and underexplored this field remains. The real opportunities may lie not only in efficiency gains but also in reshaping how professional standards, trust, and accountability are embedded into legal practice. In this way, agentic AI could become both a productivity tool and a force for raising the bar of professional conduct.

Foresight: key opportunities

Signal: Early adoption is potentially strongest in contracts, compliance, and knowledge management – areas ripe for autonomous, goal-directed AI systems.

Implications:

- Agentic AI could let small firms deploy “micro-agents” that negotiate, draft, and file autonomously, eroding scale advantages once held by global firms.
- Always-on agentic systems may monitor client risk and regulatory change in real time, shifting legal work from reactive advice to proactive prevention.
- By 2040, firms may operate as orchestration layers for swarms of legal agents, blurring the line between legal service provider, data custodian and AI platform.
- Competitive advantage may hinge on how well lawyers train, supervise, and collaborate with their agentic counterparts, creating new forms of professional hierarchy.
- If younger lawyers are more fluent in tech and AI, but not experienced in the law and complex legal matters, does this foreshadow a shift towards greater reliance on AI systems to know and produce the law?

6. How might identity, reputation and accountability evolve for lawyers in an agentic AI-integrated landscape?

Synopsis

In an agentic AI-integrated landscape, lawyers' identity is likely to shift from technical "doers" of legal work towards higher-value roles as advisers, strategists and interpreters of AI-generated outputs. Some may take on hybrid identities that combine legal, technological and governance expertise, acting as translators between disciplines. These identity shifts may challenge long-standing assumptions about legal expertise, apprenticeship and what constitutes professional judgement.

The profession's reputation faces both opportunities and risks. AI may widen access to legal information, but could also reinforce public perceptions that lawyers are costly, slow or protectionist. Survey data shows that although trust in lawyers remains relatively high in the UK, global perceptions are more mixed, and public experiences of delay or poor communication continue to erode confidence.

As AI tools become more capable, some clients may try to "play lawyer" or rely on built-in legal processes offered by technology platforms. Lawyers who focus on empathy, context and strong interpersonal skills may retain a reputational advantage in areas where AI cannot replicate human understanding. However, unregulated providers and fast-moving technology platforms may start to influence what the public expects from legal services, potentially reshaping how the profession is viewed.

Accountability is likely to become more complicated as autonomous or semi-autonomous systems make decisions that blur the boundary between human and machine responsibility. When responsibility is spread across developers, operators and users (the so-called "moral crumple zone") it creates difficult questions about who is liable and how decisions should be made transparent and well-governed. Research also shows that simply labelling something as AI-generated can reduce people's confidence in its quality, which further affects how accountability and trust are understood. In this context, lawyers have an important role in shaping clear frameworks that keep agentic AI systems open to challenge, capable of being audited and aligned with professional and ethical standards.

Group discussions emphasised that lawyers must retain ultimate accountability and actively shape how they remain involved in key decisions as AI expands across legal, compliance and business functions. Reputations may come to reflect both the lawyer and the AI systems they use.

Before addressing possible societal perspectives on lawyers' use of agentic AI, it is worth noting the profession's perception of AI itself. Research by Harasta, Novotná & Savelka (2024) explored how lawyers and law students perceived documents they believed are generated by AI versus those they believed to be human-written. Participants were asked to score the quality and correctness of documents labelled as either AI-generated or human-crafted. Unbeknownst to them, all the documents had been created by humans. The findings showed a consistent bias: documents labelled as AI-generated were rated lower, even when they were factually accurate. This reflects a broader phenomenon

known as “algorithmic aversion”. The study concluded that labelling content as AI-generated tends to create negative associations, such as assumptions of lower quality or reduced adequacy in legal contexts.²²

Surveys have found that public perception of lawyers in the UK is generally positive. In October 2024, the SRA published survey data to show that around three quarters of consumers and SMEs had both trust and confidence in legal services.²³ In the same year, the Legal Services Consumer Panel tracker survey recorded that satisfaction with the service provided by lawyers was 87%.²⁴

Experiences of using a solicitor are not always positive however. The SRA’s First Tier Complaints Report 2023, reported that out of 36,887 complaints received by law firms in England and Wales the most common complaints related to delay (20%); failure to keep informed (17%); failure to progress (12%); and failure to advise (10%).²⁵

Globally, the picture is also less positive. While personal service or personal experience might show high levels of satisfaction, there is evidence to suggest the legal profession’s general societal impact is not so highly regarded. The International Bar Association’s report on the social and economic impact of the legal profession (2024),²⁶ considered the impact of the legal profession on a variety of areas. The survey received over 7,500 responses from representatives of the general population in 16 countries. While assessments of “negative impact” were generally in single percentage points, the survey showed large sections of respondents did not know about the impact of lawyers, or thought they made no impact, perhaps leaving open the possibility that the reputation of the profession is more open to question.

Potential impact on accountability

Agentic AI may raise the stakes for legal professionals if there is a lack of clarity over where responsibility for actions lie. As one recent research paper points out (Mukherjee & Chang 2025), what happens when “accountability is misattributed to a human actor who had little real agency over the system’s behaviour”?²⁷ With the increasing use of AI in the legal profession there is also the view that AI systems, when properly used, are open and responsive to human dispute and intervention. That is humans have control over AI (“contestability”) and that can still interrupt, redirect and instruct an AI system.²⁸

These issues contribute to what Mukherjee and Chang (2025)²⁹ have labelled the “moral crumple zone” problem. A problem created by agentic AI, where accountability gets diffused among developers, operators and users. Their work underlines the need for:

- New liability frameworks to clarify who is answerable when autonomous AI decisions cause harm or legal disputes.
- Stronger transparency and explainability requirements so that AI-driven legal actions can be contested and understood.
- Adapted intellectual-property and competition law, as AI becomes both creator and market actor.
- Interdisciplinary governance that blends law, economics and technology to manage emergent AI “societies” without undermining fairness or trust.

In group discussions, members strongly emphasised that lawyers must retain ultimate accountability and cannot delegate responsibility to agentic AI. They argued that human judgement remains indispensable because legal practice involves persuasive, high-stakes, context-dependent decision-making that permits very little margin for error.

"You don't get to say the computer did it."

One participant noted that lawyers depend heavily on getting things right the first time, underscoring that current AI systems cannot yet deliver the level of refined legal judgement that the profession requires.

Some members also suggested that the legal profession may be relying on a flawed assumption: that lawyers will keep working in the same way, simply with AI added into existing processes. They warned that as AI becomes more capable, clients may increasingly bypass lawyers ("I've already asked ChatGPT"), shifting the challenge from efficiency to ensuring lawyers remain involved in key decisions. Participants with in-house experience stressed the importance of preserving the integrity and authority of the lawyer-led legal stream as AI tools spread across legal, compliance, risk and business functions. Their concern was that the role of the lawyer could be sidelined or significantly reshaped unless lawyers take a more active role in defining how they stay part of the process.

Potential impact on reputation

Most interviewees perceived the legal profession to suffer from a reputation problem: they considered that people distrust lawyers, resent their costs and feel let down by their responsiveness. For interviewees, this left the profession vulnerable to disruption by agentic AI, which may be welcomed as a cheaper and less frustrating alternative. Some interviewees noted a tension: while people resent lawyers as expensive gatekeepers, they also fear the consequences of lacking professional legal support.

"Trust in the legal profession is in many parts of society pretty low. Most people can't get hold of a lawyer. Most people think lawyers are too expensive. They think we are protectionist."

In the US, one interviewee noted that public trust is already very low, suggesting that agentic AI could improve perceptions by addressing the access-to-justice gap, with around 80% of people unable to afford civil legal services. By contrast, others pointed out that in the UK, surveys continue to show relatively high trust in lawyers, even after scandals such as the Post Office Horizon case. This highlights a disconnect: insiders see systemic failings, while public opinion is shaped more by personal experiences.

Some interviewees warned that the biggest barrier to the use of agentic AI would not be public trust but resistance within the profession itself, driven by ignorance, insecurity and fear of the unknown. They noted that the challenge for professional bodies like the Law Society will be to lead with education, reassure practitioners and help them adapt.

Some interviewees predicted a fundamental shift in lawyers' daily work. Lawyers could focus on higher-value roles: advising clients on preventing disputes, offering strategic guidance and managing the technology that does the 'grunt work'.

"You'll have to speak not only law, you'll have to speak economics, and you'll have to speak code, and you'll have to speak technology."

"I went through a law degree selection process. They look for people with great analytical skills and the ability to read a lot. And actually you're not going to need any of that, sorry to say ... You're not going to need to read a lot because a really good agentic AI tool would be able to do most of the summarising for you and draw out the most pertinent points and do the analysis. The bit that's going to be most important is probably the interpersonal relationships... interpreting that for the needs of your clients."

Some interviewees described hybrid identities and roles: lawyers as data scientists, CTOs, or translators between law and technology.

"So actually the lawyer goes from being topic specific, expert driven power base to a generic ability to translate that expertise to many different sectors... You could almost liken it to becoming more of a consultancy model than a legal expertise model."

Concerns were raised that agentic AI could weaken public perceptions of the profession. Examples were cited of organisations already receiving "very well worded but wrong" generative AI created complaints. This suggests the public may increasingly feel empowered to 'play lawyer' without expertise, risking both outcomes and professional credibility.

"If people start to think lawyers and law firms are using agents and generative AI, they can maybe think they can be a lawyer themselves... does it diminish how people see lawyers as a profession?"

The threat from legal platforms

Some interviewees predicted the greatest disruption will come from technology platforms absorbing legal processes directly into their terms of service (e.g., Upwork for IP and contracts, or rental platforms for leases and guarantees).

Such platforms standardise and automate arrangements, bypassing lawyers entirely. This could erode the profession's centrality, with lawyers' complexity and protectionism seen as barriers rather than value.

Members in the group discussions echoed this concern that if agentic AI becomes widely used and big tech companies enter the legal market, lawyers could be pushed out of parts of their own work. They considered that faster-moving vendors, client self-service tools and unregulated providers may intensify competitive pressure on regulated lawyers who must follow strict ethical and professional rules.

They were concerned about fairness and integrity: if unregulated tech providers could offer something that looks similar to legal services but without the same obligations, what happens to quality, accountability and access to justice? Members saw a tension between encouraging innovation and protecting the ethical safeguards that underpin the profession.

The opportunity to reframe professional identity

Most interviewees were optimistic that whilst AI, and particularly agentic AI, threatens lawyers' traditional identity, those who embrace governance, ethics, and new skills will remain relevant and could even strengthen their value:

- **Governance and oversight:** lawyers can position themselves as leaders in agentic AI governance, ensuring systems are compliant, accountable, and ethically deployed.
- **Ethical leadership:** embedding traditional principles (fiduciary responsibility, confidentiality, fairness) into AI-augmented practice maintains continuity with core values.
- **New skills:** learning about technology, risk management, and working in teams that mix lawyers with technical experts.

"Law firms are going to have to sell insight, foresight and certainty rather than billable hours and advice... there's opportunity to be trusted and to show what other skills are required... and to reframe what it means to be a solicitor."

In the group discussions, members emphasised professional identity and legal training as essential anchors in an era of agentic AI. They considered that lawyerly judgement, apprenticeship, close reading and critical reasoning must be preserved to maintain what makes the profession distinctive. These conversations also surfaced deeper questions about the nature of expertise, the craft of legal practice and what counts as professional knowledge.

Despite technological advances, interviewees perceived that clients would continue to value fundamentally human qualities in their legal advisers, including empathy, judgement and the ability to navigate complex interpersonal and commercial considerations that extend beyond pure legal analysis. This preference for human connection may provide a reputational anchor for lawyers who can demonstrate these capabilities alongside technical competence.

"The bit that's going to be most important is probably the interpersonal relationships... interpreting [legal knowledge] for the needs of your clients... That's all about interpersonal relationships. It's all about EQ, it's all about understanding your clients' markets more than understanding the law."

"A legal answer to a client's problem doesn't always involve them following what the legal position is. There's many commercial and emotional and all kinds of other factors... getting the correct legal bit is only part of solving the client's problem."

In the group discussions, members raised concerns about responsibility, professional identity and the boundaries of automation. They reflected on how agentic AI could challenge long-standing assumptions about legal agency, expertise and accountability. Their concerns also extended to broader societal impacts: how agentic AI might reshape professional markets, redistribute risks and benefits, affect workers and future generations, and place new pressures on regulatory and insurance frameworks. Members ultimately saw AI as a systemic force with wide-ranging ripple effects that demand coordinated and ethical governance.

Key reflections on identity, reputation and accountability

Changes to lawyers' daily roles may reduce the emphasis on traditional skills such as reading and analysis but there will be an increase in the importance of interpersonal skills, judgement, and the ability to translate complex knowledge for clients. In the age of AI the enduring value of the profession will rest in qualities AI cannot replicate such as empathy, judgement and the ability to navigate complex human and commercial contexts.

Accountability challenges mean work needs to be undertaken to develop new liability frameworks as well as approaches to transparency, explainability and governance.

Foresight: identity, reputation and accountability

Signal: By 2040, client trust and professional credibility will be increasingly mediated by AI agents, not just human lawyers.

Implications:

- Invisible AI mediation may fundamentally reshape the profession's identity: clients could interact primarily with intelligent systems, leaving human judgement perceived as optional.
- Reputation may become multi-layered, applying both to lawyers and their AI agents, "trusted agent certification" or AI track records could become standard metrics of credibility.
- Law firms may polarise along new lines: some embracing a "human-first" ethos, emphasising judgement, ethics and personal touch; others promoting "AI-optimised law," where efficiency, predictive insight and agentic autonomy define the brand.
- Ethical accountability may increasingly rely on AI itself, with specialised agentic oversight ensuring compliance, bias mitigation and transparency, potentially redefining what it means to be a responsible lawyer.

7. Conclusions

Beyond today: futures of law in an agentic age

The evidence gathered for this report shows that fully autonomous agentic AI is not yet embedded in legal practice or justice systems. Human oversight remains the rule and professional responsibility remains paramount.

While the full consequences of agentic AI remain uncertain, its development is likely inevitable. Responsible foresight—considering environmental, ethical, market, and technological risks—is essential to ensure that adoption strengthens, rather than undermines, professional standards, client trust, and sustainable practice. By anticipating these challenges now, the profession can guide AI's integration rather than react to it.

But the trajectory is clear: each year, systems take on more planning, orchestration and autonomy.

The key foresight insight is that change may come less through a sudden breakthrough than through incremental drift:

- Lawyers delegating “just one more task” to systems.
- Judges consulting AI outputs to ease backlogs.
- Firms restructuring around workflows where AI is the first point of contact.

If left unaddressed, this drift could normalise invisible substitution of professional judgement, the emergence of quasi-institutional AI actors, and a gradual shift of authority away from humans toward systems that are powerful but unaccountable. Environmental sustainability may become a decisive constraint on agentic AI adoption. As real-world energy-use data improves, firms should expect sustainability considerations to shape regulatory expectations and client due-diligence.

By 2040, the legal landscape could look radically different:

- Scenario A: Agentic AI remains confined to administration, relieving human lawyers to focus on complex judgement.
- Scenario B: Agentic AI becomes the de facto paralegal, handling the majority of routine client interactions and filings.
- Scenario C: In some jurisdictions, agentic AI systems are recognised as formal decision-makers – raising profound questions of legitimacy, sovereignty, and justice.

For the Law Society and its members, the foresight challenge is twofold:

1. To prepare the profession – equipping lawyers to work alongside, oversee, and challenge agentic AI.
2. To shape the future – ensuring that law, regulation, and governance evolve fast enough to prevent systems that act with institutional weight from escaping institutional accountability.

In short: agentic AI is not just another tool. It is a potential new actor in the legal ecosystem. The question for the next two decades is whether the legal profession will shape its role – or be shaped by it.

The role of the Law Society

Agentic AI is not yet embedded in legal practice, but its potential to reshape law is qualitatively different from generative AI. The challenge for the profession and for the Law Society is to get ahead of the curve, not to be caught reacting once systems are already acting with institutional weight.

In the group discussions, Law Society members responded to the following ideas and shared their views on the role that could be played by the Law Society concerning agentic AI.

Proactive governance: The Society can help shape standards, certification and best practices for agentic AI adoption before regulation is imposed externally, ensuring accountability, transparency and professional integrity.

Members wanted the Law Society to lead the conversation about generative and agentic AI and legal practice. They saw the Law Society as a natural convenor to pull professional consensus together.

Most agreed that the Society needed to provide clear, practical and proportionate guidance rather than advocating for new layers of regulation.

"If there's a role for the Law Society, it is in leadership ... values and integrity and ethics ... giving us those parameters"

Members wanted the Law Society to coordinate with regulators and professional bodies (ICO, SRA, insurers, Bar Council, CILEX, LSB) to avoid fragmented or conflicting rules about agentic AI. They wanted the Law Society to clarify whether existing legal rules and professional standards adequately cover generative and agentic AI use, or if there is a need for new frameworks. Members face these questions daily and need clear, practical answers to practice confidently and compliantly.

Members considered that it was important that the Law Society engages with insurers and clarifies professional indemnity expectations so that firms understand when generative (and potentially agentic) AI use affects coverage and how risk should be managed.

Some members cautioned against overstating the threat posed by agentic AI, reminding others that *"machines are not a legal entity."*

Education and clarity: Lawyers and clients struggle to define agentic AI and its unique implications. The Society can provide authoritative guidance to cut through hype, distinguish between tools and quasi-institutional systems, and explain ways in which agentic AI can and cannot practise law.

Members wanted to receive clear, practical guidance and education so that the profession can develop a shared, accurate understanding of agentic AI capabilities and limits and understand the differences between generative and agentic AI systems.

"There's still a lot of base level education that is required on different terms, how AI works, what the risks are."

They wanted to know more about what agentic AI can do in legal practice, with task-bounded guidance using real examples. Members wanted help with practical legal questions they might face in daily practice such as disclosability and client privilege, liability, insurance coverage and verification versus validation. They wanted to understand where existing principles suffice and where new regulatory clarity is needed.

Ethics infrastructure for agentic AI: Beyond rules, the Society can develop frameworks to embed ethical reflection into agentic AI use, ensuring systems augment professional judgement rather than displace it or shift moral responsibility onto algorithms.

Members broadly supported the Law Society developing practical ethics and governance infrastructure to ensure that agentic AI systems enhance rather than displace professional judgement. They wanted the Law Society to set clear expectations about human oversight and provide guidance, standards and guardrails. They emphasised that such frameworks should be practical and operational rather than abstract or vague.

Foresight scenario testing: The Society can explore futures in which agentic AI operates anywhere from back-office support to quasi-institutional actors, stress-testing the profession, the justice system and regulatory frameworks against radical AI-driven scenarios.

Members wanted the Law Society to maintain foresight monitoring and future-preparedness, while prioritising immediate, usable support for practitioners – especially in the transition towards future business environments.

There was less appetite from the members consulted for the following:

Legal literacy for AI outputs: Lawyers will need to understand, scrutinise and validate agentic AI-generated contracts, advice, and litigation strategies. The Society can support with training to ensure outputs are defensible and aligned with legal principles.

There was a mixed response to the idea of the Law Society providing training. Some in-house members in particular identified that they would not look to the Law Society for this, preferring to access training from vendors or their panel law firm providers.

Preserving professional trust: The Society can guide lawyers in managing relationships with clients when agentic AI increasingly mediates advice, negotiation and drafting, protecting human reputation and accountability while allowing AI to augment work.

There was no discussion about this idea.

Bibliography

Anderson, H., Reem, N. & Susas, J., 2025. Automated Decision Making Emerges as an Early Target of State AI Regulation. White & Case LLP (Insight Alert), 7 March. Available at: <https://www.whitecase.com/insight-alert/automated-decision-making-emerges-early-target-state-ai-regulation>

American Bar Association, 2024. Formal Opinion 512: Lawyers' Use of Generative Artificial Intelligence Tools. ABA Standing Committee on Ethics and Professional Responsibility. Available at: <https://www.lawnext.com/wp-content/uploads/2024/07/aba-formal-opinion-512.pdf>

Bar Council, 2024. Considerations when using ChatGPT and generative artificial intelligence software based on large language models. Information Technology Panel, 30 January. Available at: <https://www.barcouncilethics.co.uk/documents/considerations-when-using-chatgpt-and-generative-ai-software-based-on-large-language-models>

BCLP, 2024. California turns to the use of AI in healthcare. Bryan Cave Leighton Paisner LLP. Available at: <https://www.bclplaw.com/en-US/events-insights-news/california-turns-to-the-use-of-ai-in-healthcare.html>

Brookings, 2025. How different states are approaching AI. Available at: <https://www.brookings.edu/articles/how-different-states-are-approaching-ai/>

CAC (Cyberspace Administration of China), 2023. Interim Measures for the Management of Generative Artificial Intelligence Services. Available at: <https://www.chinalawtranslate.com/en/generative-ai-interim/>

Casper, S., Bailey, L., Hunter, R., Ezell, C., Cabalé, E., Gerovitch, M., Slocum, S., Wei, K., Jurković, N., Khan, A., Christoffersen, P., Ozisik, A., Trivedi, R., Hadfield-Menell, D. & Kolt, N., 2025. The AI Agent Index. arXiv preprint arXiv:2502.01635. Available at: <https://arxiv.org/pdf/2502.01635>; Index available at: <https://aiagentindex.mit.edu>

Chambers & Partners, 2025. Artificial intelligence 2025 - China. Global Practice Guide. Available at: <https://practiceguides.chambers.com/practice-guides/artificial-intelligence-2025/china>

Chan, A., Salganik, R., Markelius, A., Pang, C., Rajkumar, N., Krashenninnikov, D., Langosco, L., He, Z., Duan, Y., Carroll, M. & Lin, M., 2023. Harms from increasingly agentic algorithmic systems. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, pp. 651-666.

College of Policing, 2025. Responsible AI Checklist for Policing. Available at: <https://library.college.police.uk/docs/NPCC/Responsible-AI-checklist-for-policing-2025.pdf>

Competition and Markets Authority, 2024. AI Foundation Models: Update paper. GOV.UK. Available at: <https://www.gov.uk/cma-cases/ai-foundation-models-initial-review>

Crown Prosecution Service, 2025. How the CPS will use Artificial Intelligence. GOV.UK. Available at: <https://www.cps.gov.uk/how-cps-will-use-artificial-intelligence>

Department for Science, Innovation and Technology & AI Safety Institute, 2025.

International AI Safety Report 2025. GOV.UK. Available at:

<https://www.gov.uk/government/publications/international-ai-safety-report-2025>

Digital Regulation Cooperation Forum, 2024a. AI and Digital Hub – DRCF. Available at:

<https://www.drcf.org.uk/ai-and-digital-hub>

Digital Regulation Cooperation Forum, 2024b. AI and Digital Hub – DRCF [pilot].

Available at: <https://www.drcf.org.uk/ai-and-digital-hub>

Emory University School of Law, n.d. AI and the Future of Work. Available at:

<https://law.emory.edu/centers-and-programs/ai-future-of-work.html>

Estonian Information System Authority (RIA), 2024. Risks and controls for artificial intelligence and machine learning systems. Tallinn: RIA. Available at:

<https://www.ria.ee/sites/default/files/documents/2024-05/Risks-and-controls-for-artificial-intelligence-and-machine-learning-systems.pdf>

Estonian Information System Authority (RIA), 2025. Bürokratt: interoperable network of public-sector virtual assistants. Available at: <https://www.ria.ee/en/state-information-system/personal-services/burokratt>

e-Estonia, 2021. National AI strategy—KrattAI factsheet. Available at: <https://e-estonia.com/new-e-estonia-factsheet-national-ai-kratt-strategy/>

e-Residency of Estonia, 2025. e-Residency of Estonia – home page. Available at:

<https://www.e-resident.gov.ee>

European Commission, 2023. Digital public services based on open source: Case study—Bürokratt. Available at: <https://interoperable-europe.ec.europa.eu/collection/open-source-observatory-osor/document/digital-public-services-based-open-source-case-study-burokratt>

European Commission, AI Watch, 2019. Estonia AI strategy report. Available at:

https://ai-watch.ec.europa.eu/countries/estonia/estonia-ai-strategy-report_en

European Commission, AI Watch, 2025. Estonia AI strategy report. Available at:

https://ai-watch.ec.europa.eu/countries/estonia/estonia-ai-strategy-report_en

Executive Office of the President, 2023. Executive Order 14110: Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. *Federal Register*, 1 November, 88 FR 75191–75226. Available at:

<https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>

Federal Trade Commission, 2025. AI and the Risk of Consumer Harm. Tech at FTC, 3 January. Available at:

https://data.aclum.org/storage/2025/01/FTC_www_ftc_gov_policy_advocacy-research_tech-at-ftc_2025_01_ai-risk-consumer-harm.pdf

Fuller, S., Woodhall, C. & Walker, I., 2024. Making an impact. International Bar Association, 31 July. Available at: <https://www.ibanet.org/making-an-impact>

Geistfeld, M., 2024. Principles of the Law, Civil Liability for Artificial Intelligence. The American Law Institute. Available at: <https://www.ali.org/project/principles-law-civil-liability-artificial-intelligence>

Gordon, D. & Nouwens, M., 2022. Introduction The Digital Silk Road: China's technological rise and the geopolitics of cyberspace. *International Institute for Strategic Studies (IISS), Adelphi Online Analysis*, 6 December. Available at: <https://www.iiss.org/online-analysis/online-analysis/2022/12/digital-silk-road-introduction>

Harasta, J., Novotná, T. & Savelka, J., 2024. It Cannot Be Right If It Was Written by AI: On Lawyers' Preferences of Documents Perceived as Authored by an LLM vs a Human. arXiv preprint arXiv:2407.06798. Available at: <https://arxiv.org/abs/2407.06798>

Hepburn, J.D., 2024. Florida Bar Provides Guidance on Lawyers' Use of Generative AI. McLane Middleton Insights, 20 March. Available at: <https://www.mclane.com/insights/florida-bar-provides-guidance-on-lawyers-use-of-generative-ai/>

HM Government, 1976. Solicitors (Northern Ireland) Order 1976, SI 1976/582 (N.I. 12). Available at: <https://www.legislation.gov.uk/nisi/1976/582>

HM Government, 2007. Legal Services Act 2007, c.29. Available at: <https://www.legislation.gov.uk/ukpga/2007/29>

Jasnow, D., 2025. New Lawsuits Targeting Personalized AI Chatbots Highlight Need for AI Quality Assurance and Safety Standards. *National Law Review*, 6 January. Available at: <https://natlawreview.com/article/new-lawsuits-targeting-personalized-ai-chatbots-highlight-need-ai-quality-assurance-and-safety-standards>

Judiciary of England and Wales, 2025. Artificial Intelligence (AI): Guidance for Judicial Office Holders. 14 April. Courts and Tribunals Judiciary. Available at: <https://www.judiciary.uk/wp-content/uploads/2025/04/Refreshed-AI-Guidance-published-version.pdf>

Kaspar Korjus, 2018. Estonian president Kersti Kaljulaid reveals the future direction of e-Residency. *e-Residency blog*, 18 December. Available at: <https://www.e-resident.gov.ee/blog/posts/estonian-president-kersti-kaljulaid-reveals-the-future-direction-of-e-residency-2/>

KPMG, 2025. State series: AI legislation regulatory alert. Available at: <https://kpmg.com/kpmg-us/content/dam/kpmg/pdf/2025/state-series-ai-legislation-reg-alert.pdf>

Loring, J.M., 2025. When AI Acts Independently: Legal Considerations for Agentic AI Systems. *The National Law Review*, 10 June. Available at: <https://www.natlawreview.com/article/when-ai-acts-independently-legal-considerations-agentic-ai-systems>

Legal Services Board, 2025. Statement on High Court decision in Mazur v Charles Russell Speechlys LLP. Available at: <https://legalservicesboard.org.uk/news/lsb-statement-high-court-decision-in-mazur-13-october-2025>

Mansi, G. & Riedl, M., 2024. Recognizing lawyers as AI creators and intermediaries in contestability. arXiv [Preprint], arXiv:2409.17626. doi: 10.48550/arXiv.2409.17626

MultiState, 2025. AI legislative updates, vol. 71. Available at: <https://www.multistate.ai/updates/vol-71>

Mukherjee, A. & Chang, H.H., 2025. Agentic AI: Autonomy, Accountability, and the Algorithmic Society. arXiv preprint arXiv:2502.00289. Available at: <https://arxiv.org/abs/2502.00289>

Mukherjee, A. & Chang, T., 2025. Agentic AI: Autonomy, Accountability, and the Algorithmic Society. arXiv preprint arXiv:2502.00289. Available at: <https://arxiv.org/abs/2502.00289>

National New Generation Artificial Intelligence Governance Specialist Committee, 2021. Ethical Norms for New Generation Artificial Intelligence (trans. Centre for Security and Emerging Technology). Available at: https://cset.georgetown.edu/wp-content/uploads/t0400_AI_ethical_norms_EN.pdf

National Police Chiefs' Council, 2023. Covenant for using Artificial Intelligence in Policing. Police Science and Technology. Available at: <https://science.police.uk/delivery/resources/covenant-for-using-artificial-intelligence-ai-in-policing/>

North Carolina Central University, n.d. Course Catalog: [Course details page]. NCCU eCatalog. Available at: https://ecatalog.nccu.edu/preview_course_nopop.php?catoid=26&coid=44237

Northern Ireland Assembly Research and Information Service, 2024. Artificial Intelligence regulation, use and innovation in the United Kingdom: a broad overview (Paper 11/24, NIAR 45-24). Available at: <https://www.niassembly.gov.uk/globalassets/documents/raise/publications/2022-2027/2024/economy/1124.pdf>

Northern Ireland Executive, 2025. New AI and Digital Office will drive innovation and transform services – O'Neill and Little-Pengelly. Available at: <https://www.executiveoffice-ni.gov.uk/news/new-ai-and-digital-office-will-drive-innovation-and-transform-services-oneill-and-little-pengelly>

O'Grady, C.G. & O'Grady, C., 2025. Agentic Workflows in the Practice of Law—AI Agents as Ethics Counsel. *Arizona Legal Studies Discussion Paper* No. 25-03; 38 *Georgetown Journal of Legal Ethics*, pp.247–. Available at: <https://ssrn.com/abstract=5113810>

Sapkota, R., Roumeliotis, K.I. & Karkee, M., 2025. AI Agents vs. Agentic AI: A Conceptual Taxonomy, Applications and Challenges. *Information Fusion*. Available at: <https://arxiv.org/abs/2505.10468>

Salinski, D., 2025. In America: How Mobley v Workday Inc. Is Quietly Reshaping Legal and Legislative AI Accountability in the United States. *Modern Treatise*, 29 June. Available at: <https://www.moderntreatise.com/the-americas/2025/6/29/in-america-how-mobley-v-workday-inc-is-quietly-reshaping-legal-and-legislative-ai-accountability-in-the-united-states>

Scottish Government, 2021. Scotland's artificial intelligence strategy: Trustworthy, ethical and inclusive. gov.scot. Available at: <https://www.gov.scot/publications/scotlands-ai-strategy-trustworthy-ethical-inclusive/>

Scottish Parliament, 2010. Legal Services (Scotland) Act 2010, asp 16. Available at: <https://www.legislation.gov.uk/asp/2010/16/contents>

Seyfarth, 2024. Artificial intelligence legal roundup: Colorado postpones implementation of AI law. Seyfarth Shaw LLP. Available at: <https://www.seyfarth.com/news-insights/artificial-intelligence-legal-roundup-colorado-postpones-implementation-of-ai-law-as-california-finalizes-new-employment-discrimination-regulations-and-illinois-disclosure-law-set-to-take-effect.html>

Social Media Victims Law Centre, 2025. AI Chatbot Lawsuit Clears Key Legal Hurdle in Landmark Case. Blog, 23 May. Available at: <https://socialmediavictims.org/blog/character-ai-lawsuit-moves-forward>

Solicitors Regulation Authority, 2019. Code of Conduct for Solicitors, RELs and RFLs. SRA. Available at: <https://www.sra.org.uk/solicitors/standards-regulations/code-conduct-solicitors/> (Accessed: 1 October 2025)

Solicitors Regulation Authority, 2022. Compliance tips for solicitors regarding the use of AI and technology. Available at: <https://www.sra.org.uk/solicitors/resources/innovate/compliance-tips-for-solicitors/>

Solicitors Regulation Authority, 2023. Risk Outlook: The use of artificial intelligence in the legal market. SRA.

Solicitors Regulation Authority, 2024a. Research shows trust in legal services high: Strategy benchmarking survey 2024. Available at: <https://www.sra.org.uk/sra/news/press/2024-press-releases/strategy-benchmarking-survey-2024/>

Solicitors Regulation Authority, 2024b. First Tier Complaints Report 2023. Available at: <https://www.sra.org.uk/sra/research-publications/first-tier-complaints-2023/>

Surden, H., 2024. Artificial intelligence and the practice of law: Risks and responsibilities. *Yale Journal of Law and Technology*, 26(3), pp.101-138

Tabassi, E., 2023. Artificial Intelligence Risk Management Framework (AI RMF 1.0). National Institute of Standards and Technology, Gaithersburg, MD. Available at: <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>

The Law Society, 2024. Generative AI: the essentials. Available at: <https://www.lawsociety.org.uk/topics/ai-and-lawtech/generative-ai-the-essentials>

The Law Society, 2025. Mazur v Charles Russell Speechlys: what it means for litigators. Available at: <https://www.lawsociety.org.uk/topics/civil-litigation/mazur-v-charles-russell-speechlys-litigators>

The Law Society Gazette, 2024. Satisfaction levels and trust in lawyers continue to rise. Available at: <https://www.lawgazette.co.uk/news/satisfaction-levels-and-trust-in-lawyers-continue-to-rise/5120359.article>

Trump, D.J., 2025. Removing Barriers to American Leadership in Artificial Intelligence (Executive Order 14179). *Federal Register*, 31 January, 90 FR 8741–8742. Available at: <https://www.federalregister.gov/documents/2025/01/31/2025-02172/removing-barriers-to-american-leadership-in-artificial-intelligence>

TrustArc, 2025. Utah's AI Policy Act Is Here. Is Your AI Ready to Confess? Available at: <https://trustarc.com/resource/utah-ai-policy-act>

UK Government, 2023. Safety Testing: Chair's Statement of Session Outcomes, 2 November. Available at: <https://assets.publishing.service.gov.uk/media/6544ec4259b9f5001385a220/aiss-statement-on-safety-testing-outcomes.pdf>

UK Government, 2025a. The Bletchley Declaration by countries attending the AI Safety Summit, 1–2 November 2023. GOV.UK. Available at: <https://www.gov.uk/government/publications/ai-safety-summit-2023-the-bletchley-declaration/the-bletchley-declaration-by-countries-attending-the-ai-safety-summit-1-2-november-2023>

UK Government, 2025b. A Pro-Innovation Approach to AI Regulation: Amended Government Response. Available at: <https://assets.publishing.service.gov.uk/media/65c1e399c43191000d1a45f4/a-pro-innovation-approach-to-ai-regulation-amended-governement-response-web-ready.pdf>

UK Parliament, 2025. Artificial Intelligence (Regulation) Bill [HL]. Available at: <https://bills.parliament.uk/bills/3942>

Utah Legislature, 2024a. SB 149 Artificial Intelligence Policy Act. Utah State Legislature. Available at: <https://le.utah.gov/~2024/bills/sbillint/SB0149.htm>

Utah State Legislature, 2024b. Senate Bill 149: Generative Artificial Intelligence Disclosure Requirements, Chapter 2, Section 12. Available at: <https://le.utah.gov/~2024/bills/static/SB0149.html#13-2-12>

Appendix: Details of AI regulatory context in UK, US, Estonia and China

1. AI regulatory context in the United Kingdom

The UK does not have a single, comprehensive AI Act. Instead, government is following a pro-innovation model, working with existing regulators such as the Information Commissioner's Office, the Competition and Markets Authority and the Financial Conduct Authority.

UK government convened international AI Safety Summit

In 2023 the UK government convened the AI Safety Summit which was attended by 29 countries. Attendees agreed to the Bletchley Declaration (UK Government, 2023)³⁰ on AI safety, a landmark agreement recognising a shared consensus on the opportunities and risks of AI and the need for collaborative action on frontier AI safety. A number of countries together with frontier AI developers agreed to state-led safety evaluations of next-generation models before deployment, including via partnerships with AI Safety Institutes (UK Government, 2023).³¹

UK sets out context-based framework with existing regulators

In February 2024 the government published its formal response to the AI White Paper,³² confirming a context-based framework with regulators applying existing laws with AI-specific guidance, rather than creating one sweeping statute. Protections are already provided by regulators for example article 22 of UK GDPR, grants individuals the right not to be subject to a decision based solely on automated processing that produces legal or significant effects.

The formal response to the AI White Paper also explicitly explores the idea of new responsibilities for developers of "highly capable general-purpose AI" (i.e., the systems most likely to exhibit agentic behaviour) and sets out an AI regulation roadmap. The *AI Opportunities Action Plan* (UK Government, 2025) commits to the following (among other proposals):

- Funding to scale up regulators' AI expertise and capabilities, particularly where urgent gaps are identified.
- Encouraging sponsor departments to include enabling safe AI innovation in their strategic guidance to regulators.
- Working with regulators to accelerate AI adoption in priority sectors and deploying pro-innovation tools such as regulatory sandboxes.
- Requiring regulators to report annually on how they have supported AI-driven innovation and growth in their domain, including metrics such as timelines for guidance, licensing decisions and AI-related resource allocation.

The work internationally continued in 2025 with publication of the International AI Safety Report³³ (Department for Science, Innovation and Technology & AI Safety Institute, 2025). This report focuses on risks and technical approaches to risk management and will inform ongoing international discussions.

Artificial Intelligence (Regulation) Bill re-introduced for discussion

The Artificial Intelligence (Regulation) Bill was reintroduced in the House of Lords on 4 March 2025 by Lord Holmes as a Private Member's Bill. It proposes a regulatory

structure distinct from the current government model. Among its key features, it would designate an AI Officer, support regulatory sandboxes and provide a regulatory (if enacted) establish an AI Authority, require firms that develop, deploy or use AI to “scaffolding” via delegated regulations to impose obligations on more autonomous or agentic AI systems.³⁴

AI safeguards in the UK

The UK approach to AI safeguards emphasises responsible use underpinned by human oversight, transparency, accountability and security.

Courts and Tribunals Judiciary Guidance on AI

In April 2025 the Courts and Tribunals Judiciary issued updated Judicial Guidance on AI. The guidance reinforces the importance of developing an understanding of AI and its applications, the importance of upholding confidentiality and privacy, that AI can be inaccurate and biased, the importance of maintaining security, underlining who takes responsibility and noting that provided AI is used responsibly, there is no reason why a legal representative ought to refer to its use but it may be necessary at times to remind individual lawyers of their obligations to verify accuracy. The guidance also provides information on Microsoft’s “Copilot Chat” being available to judicial office holders.³⁵

Crown Prosecution Service Guidance on AI

In June 2025 the Crown Prosecution Service set out principles for its use of AI:

- AI will be overseen by humans, with human participation, output verification, and senior responsible officers.
- The CPS will promote transparency in its use of AI, communicating clearly to staff, partners, the public, and affected parties.
- It will continuously evaluate and improve AI systems, monitoring risks, soliciting feedback, and addressing harms.
- CPS will collaborate with Criminal Justice System partners and wider government on AI approaches, sharing knowledge and aligning practices.
- AI deployment must have a clear purpose and measurable impact, with processes to articulate tool functionality and justify use.
- CPS will focus on using AI securely and responsibly, consistent with its AI Data Security and Ethics Principles, with accountability, oversight, and governance.³⁶

National Police Chiefs’ Council (NPCC) Covenant for using AI

The National Police Chiefs’ Council endorsed a *Covenant for Using Artificial Intelligence in Policing* in September 2023,³⁷ committing all forces to use AI in ways that are lawful, transparent, explainable, responsible, accountable and robust. The NPCC’s more recent strategic planning and the College of Policing’s *Responsible AI Checklist* (2025) reference and build upon those principles, grounding their guidance in the Covenant’s framework.³⁸

Law Society of England & Wales advice on AI

The Law Society of England & Wales published *Generative AI: The Essentials* (2024), a guide for the legal profession that outlines both the opportunities and risks of generative AI covering issues such as data protection and privacy, misuse by malicious actors and the ethical considerations legal professionals should bear in mind.³⁹

Bar Council guidance on generative AI

The Bar Council of England and Wales published *Considerations when using ChatGPT and generative artificial intelligence software based on large language models (2024)*⁴⁰ which provides information on the key risks of LLMs and considers a range of professional considerations. It warns that practitioners should recognise the constraints and challenges presently embedded in the generative AI LLM software.

UK regulator activity on AI

Solicitors Regulation Authority (SRA) guidance

In England and Wales, the Solicitors Regulation Authority (SRA) applies its *Code of Conduct for Solicitors* (2019).⁴¹ This requires solicitors to act competently, in the best interests of each client and to provide a proper standard of service. These obligations extend to the use of AI, even though AI is not mentioned specifically in the Code. Solicitors must also maintain their competence and keep their professional knowledge and skills up to date.

In 2022, the SRA published *Compliance Tips for solicitors regarding the use of AI and technology*, answering practical questions about lawtech adoption.⁴² Its 2023 *Risk Outlook - The use of artificial intelligence in the legal market* examines opportunities and risks in more detail.⁴³ The report highlights specific risks when legal professionals rely on AI, including over-reliance on outputs, the need to disclose use of such tools to clients where relevant, and compliance with data protection requirements. It suggests that solicitors must manage these risks so that any AI systems used operate consistently with professional duties of safety, accuracy and lawfulness. The SRA also notes that its regulation focuses on the outcomes that firms achieve, not the particular tools used to deliver legal services.

Information Commissioner's Office (ICO) guidance

The ICO's *"Guidance on AI and Data Protection"* (updated March 2023)⁴⁴ and its companion *Explaining Decisions Made with AI*⁴⁵ remain cornerstone guidance for applying data protection law to AI systems, and the March 2023 update was made in response to industry requests for clarity on fairness in AI.

Competition and Markets Authority (CMA)

Meanwhile, the Competition and Markets Authority (CMA) has pursued work on foundation models (with a technical update published in April 2024),⁴⁶ identifying competition risks and setting principles to guide foundation model markets (e.g. fairness, access, transparency). Although the CMA's work does not regulate agentic AI directly, its principles for foundation model markets shape how these powerful base models are developed, shared and used, which in turn affects the conditions under which more agentic systems can be built.

Digital Regulation Cooperation Forum (DRCF)

The Digital Regulation Cooperation Forum (DRCF) was established in July 2020, bringing together four UK regulators (CMA, ICO, Ofcom, and later the FCA) to foster greater cooperation in digital regulation.⁴⁷ In April 2024, the DRCF launched its AI and Digital Hub, a 12-month pilot offering free informal advice to innovators with regulatory questions spanning more than one regulator's remit. The pilot is now closed to new queries, and the insights gained are being used to help the DRCF and its regulators

design better regulatory support tools to help innovators bring products responsibly to market.⁴⁸

Scotland

Scotland has its own government-led AI strategy (2021-2025) which emphasises trustworthy, ethical and people-focused AI.⁴⁹ Regulation in devolved areas (such as justice, health and local government) is shaped by Scottish institutions and sector regulators, while UK-wide regulators like the ICO, CMA and FCA apply only in reserved matters such as data protection, competition and financial services. As a result, Scotland's AI governance is a blend of UK-level oversight and Scotland-specific strategy rather than a unified statutory regime.

Northern Ireland

Northern Ireland has no dedicated AI regulatory statute and as in Scotland, relies on a mix of UK-wide regulators for reserved areas and NI-specific regulators for devolved sectors.⁵⁰ Northern Ireland's regional framework is more emergent and less developed than in Scotland. In June 2025, the NI Executive established a new Office of AI and Digital as part of its Programme for Government 2024-2027, with a mandate to support the creation of "a new national AI strategy and action plan" for Northern Ireland.⁵¹

Foresight: regulation in the United Kingdom

Regulatory direction signals

- The government's AI Opportunities Action Plan and White Paper response show a gradual shift from voluntary coordination to formalised oversight, signalling that a statutory regime may eventually gain traction.
- Creation of the AI Safety Institute and leadership of international coordination indicate growing UK ambition to be a global safety and governance hub.
- Increased regulator resourcing and expectations for annual AI reports signal a trend toward measurable regulatory performance and accountability.

Cultural and institutional signals

- Judicial, CPS and policing guidance on AI shows early mainstreaming of professional ethics and responsibility across the justice system.
- Regulators like the SRA and CMA are building AI-specific expertise - a signal of regulatory professionalisation.
- A consistent emphasis on "pro-innovation but safe" policy indicates regulatory pragmatism rather than ideological opposition to AI.

Weak signals for 2030-2035

- Growing policy discussion of "highly capable general-purpose AI" suggests future differentiation between standard and agentic systems.
- Movement toward mandatory pre-deployment testing and safety certification could become visible before 2030.

2. AI regulatory context in the United States

The US lacks a comprehensive federal AI law, leaving a patchwork of state-level measures focused mainly on transparency, consumer protection, and oversight of high-risk automated decision systems that create a fragmented but evolving regulatory landscape.

Federal context

There is no comprehensive federal legislation specifically addressing AI (or agentic AI) in the US.⁵²

State context

In 2025, all US states introduced some form of AI-related legislation, with over 1,000 bills considered nationally.⁵³ However, as of August 2025 only 118 measures (roughly 11%) were enacted into law across 38 states, creating a fragmented and evolving regulatory landscape (MultiState, 2025).⁵⁴ These state-level initiatives reflect growing concern over the risks of artificial intelligence and commonly focus on transparency, consumer protection, and accountability. A notable trend across jurisdictions is the introduction of disclosure requirements, mandating that individuals be informed when AI is in use, as well as oversight provisions targeting “high-risk” automated decision systems that significantly affect people’s rights or opportunities (Brookings, 2025; NCSL, 2025).

Transparency, accountability and consumer protection

A key aspect in AI-related legislation is transparency about what AI is being used for, ensuring disclosure of when and how AI is employed. Utah was the first state to take legislative action in this area.⁵⁵ The *Utah Artificial Intelligence Act (SB149)*⁵⁶ requires disclosure of generative AI use in certain consumer-facing interactions and holds businesses accountable under consumer protection laws if such use results in violations. Licensed professionals in regulated occupations, including health professionals, must also provide advance disclosure when using generative AI (Utah Legislature, 2024).⁵⁷ State approaches vary in scope: California’s AB 3030 obliges health care providers to include disclaimers when generative AI is used in clinical communications and to give patients clear routes to human follow-up (BCLP, 2024).⁵⁸ Colorado’s SB 24-205 imposes impact assessment requirements, mandates safeguards against algorithmic discrimination and grants consumers rights over “high-risk” AI systems, however its implementation has been delayed until June 2026 (Seyfarth, 2024).⁵⁹ Other states are experimenting with narrower applications, such as professional liability frameworks or domain-specific disclosures but the core policy themes across jurisdictions remain transparency, accountability and consumer protection (Brookings, 2025).⁶⁰

“High-risk” AI interactions

“High-risk” is currently a common phrase used in AI-related legislation but there is no agreed definition between states.⁶¹ “High-risk” AI interactions have often been considered so, because of AI’s capacity for automated decision making, so this is as an overlapping area.

Automated decision making

California's *Assembly Bill 2885 (AB 2885)*⁶² introduces a statutory framework requiring state agencies to define, inventory, and oversee high-risk automated decision systems (ADS). The bill amends multiple state codes to embed a consistent definition of ADS and directs the Department of Technology (in coordination with interagency bodies) to maintain a comprehensive inventory of such systems used or procured by state agencies. Commentators suggest the law is intended to promote public accountability and guard against harmful or biased AI outcomes by establishing transparency and oversight mechanisms.⁶³

It has been asserted that the aim of AB 2885 is to ensure public accountability and legal protection against 'AI-induced biased outcomes and other potentially harmful impacts' for Californian citizens.⁶⁴

AI safeguards in the US

AI safeguards in the US remain fragmented, with general laws and standards applying alongside emerging professional guidance, most notably the American Bar Association's Formal Opinion 512 and various state bar initiatives which emphasise competence, confidentiality, transparency, supervision and accountability in lawyers' use of AI.

Federal safeguards

Use of agentic AI in the US is still developing, and consequently ethical and practical guidance is still emerging. At a national level Executive Order 14110 *Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence* was signed by President Biden on 30 October 2023. The order directed over 50 federal agencies to ensure AI systems were developed and used responsibly,⁶⁵ although it did not specifically reference AI systems with agency or autonomy. It was, however, surpassed by Executive Order 14179 *Removing Barriers to American Leadership in Artificial Intelligence* on 23 January 2025 by President Trump.⁶⁶ The Federal Trade Commission (FTC) has issued several directives concerning AI technologies, focusing on consumer protection and competition, and how it must comply with existing laws prohibiting unfair or deceptive practices, but has not specifically addressed agentic AI.⁶⁷

There are, however, some relevant frameworks, such as the National Institute of Standards and Technology's (NIST) *Artificial Intelligence Risk Management Framework*. This was released in July 2024 and gives guidance on explainability and interpretability to manage the risk of using AI.⁶⁸ The framework does not explicitly use the term agentic AI but does refer to the risks and governance of autonomous AI systems.

Legal sector safeguards

The American Law Institute (ALI) launched *Principles of the Law, Civil Liability for Artificial Intelligence* in October 2024. The project aims to develop a set of principles based on existing tort doctrines to assign responsibility for harm caused by AI systems.

Agentic AI is not explicitly mentioned, but the project addresses issues relevant to autonomous systems.⁶⁹

The American Bar Association addressed professional issues surrounding the use of generative AI in *Formal Opinion 512*⁷⁰. This opinion applies the Model Rules of Professional Conduct to the emerging challenges posed by generative AI and provides a foundation that many state bar associations may build upon. Key recommendations include:

- **Competence:** Lawyers must maintain the legal knowledge, skill, and thorough preparation required under Rule 1.1, which includes having sufficient technical understanding to assess whether an AI tool is appropriate, and independently reviewing or verifying AI outputs.
- **Confidentiality:** Lawyers must safeguard all client-related information (Rule 1.6), assess the risks of disclosing data to AI systems, and obtain informed consent where appropriate.
- **Communication:** Under Rule 1.4, lawyers must keep clients reasonably informed, which in the AI context may include explaining its use, especially if it affects decisions, costs, or outcomes.
- **Fees:** Fees must remain reasonable under Rule 1.5. Lawyers may charge for reviewing and finalising AI-generated work but not for time spent learning to use an AI tool. Where AI increases efficiency, lawyers should consider this when setting fixed or flat fees, since charging the same amount for substantially less work could be unreasonable.
- **Candor and supervision:** Lawyers must also comply with duties of candor toward tribunals (Rule 3.3) and supervisory responsibilities (Rules 5.1 and 5.3) to ensure proper oversight of both human and technological assistance.

Most, but not all, state bar associations have developed or are developing guidance to help lawyers understand their ethical obligations with regard to AI. The guidance produced varies in extent, from comprehensive frameworks to preliminary comment, and these exist within rules of professional conduct based upon the American Bar Association *Model Rules of Professional Conduct*.⁷¹ States with well-developed guidance include California, Florida and New York State.

Foresight: regulation in the United States

Regulatory direction signals

- Fragmentation of over 1,000 state-level AI bills (with 118 enacted) indicates a bottom-up federalism trend – an environment of experimentation where effective models can spread.
- The *NIST AI Risk Management Framework* and *Executive Orders 14110* and *14179* suggest eventual federal convergence on baseline safety and transparency standards, even without a single AI Act.
- The Federal Trade Commission's active stance on deceptive or harmful AI practices points to a future consumer protection route for AI accountability.

Market and institutional signals

- Rapid uptake of AI in regulated professions (law, health, finance) is prompting self-regulatory standards, particularly around competence, confidentiality and liability.
- Private insurers, auditors and risk-assurance bodies are starting to fill regulatory gaps – a market-led governance signal.
- Increasing pressure from industry coalitions for federal harmonisation may foreshadow a coordinated national framework by early 2030s.

Weak signals for 2030-2035

- "High-risk" AI categories and automated decision systems (ADS) legislation in states like California could evolve into national agentic AI oversight models.
- Federal-state collaboration (e.g. sandbox partnerships) may create multi-tiered governance, mirroring environmental or data protection regimes.

3. AI regulatory context in Estonia

Under the EU AI Act and GDPR, AI systems used in legal contexts are subject to obligations around transparency, documentation, human oversight and safeguards on automated decision-making. These frameworks do not impose special rules for agentic AI specifically, but as legal applications would fall within high-risk categories or involve personal data, strict compliance requirements would apply.

Laws which govern AI usage

Both the EU AI Act and the GDPR shape how agentic AI may be deployed in legal and administrative contexts in Estonia. Under the AI Act, AI systems that support or influence public-sector decision-making are generally classified as high-risk, triggering obligations around risk management, documentation, transparency, logging, and meaningful human oversight. Systems used in legal services may fall into high-risk categories where their outputs can affect individuals' rights or legal status.

The GDPR adds further constraints, including limits on solely automated decision-making (Article 22), enhanced transparency duties, data minimisation requirements,

and mandatory Data Protection Impact Assessments for high-risk processing of personal or special-category data.

Estonia has not enacted a standalone AI law, instead relying on amendments to existing legislation such as the Administrative Procedure Act, which governs how automated administrative decisions may be carried out, ensuring human involvement, explainability, and appeal rights (European Commission AI Watch 2025).⁷²

Together, the AI Act and GDPR create a dual regulatory regime characterised by legally required oversight, auditability and accountability. This ensures that responsibility remains with human professionals or institutions rather than AI systems themselves. In practice, the framework requires that:

- Agentic AI used in legal or administrative contexts is subject to strong governance and compliance mechanisms
- Clients and affected persons receive clear information about the role of AI and retain recourse rights
- Human professionals retain ultimate decision-making authority, with the ability to question or override AI outputs.

Estonia e-resident programme

Estonia's e-Residency programme (e-Residency of Estonia 2025), which enables global citizens to obtain a government-issued digital ID and establish businesses in Estonia remotely, has grown to over 130,800 e-residents from over 170 countries since its launch in 2014.⁷³ In 2018, during a review of the e-Residency programme, President Kaljulaid reportedly proposed granting certain rights to autonomous AI systems, although no subsequent legislation appears to have formalised this (Korjus 2024).⁷⁴

KrattAI

One of Estonia's most significant public-sector AI initiatives is KrattAI, the national framework for integrating artificial intelligence into government services. This strategic vision is implemented through Bürokratt, a virtual-assistant platform that enables citizens to access a wide range of public services through a single conversational interface (RIA 2025).⁷⁵ The name "Kratt" derives from Estonian folklore, referring to a mythical helper that performs tasks for its owner, and is used to signal that AI is intended to support, rather than replace, human decision-making (e-Estonia 2021).⁷⁶ According to the Estonian Information System Authority, Bürokratt functions as an access layer that connects users to existing state information systems, rather than autonomously carrying out administrative acts (RIA 2025). Its design principles stress interoperability, transparency, data protection and user control, including the option to escalate interactions to a human official when necessary (European Commission 2023).⁷⁷

AI safeguards in Estonia

Estonia's national AI strategy (Estonia MEAC 2021)⁷⁸ embeds plans for regulatory sandboxes, procurement reforms, and institutional capacity-building (e.g. data-steward or Chief Data Officer roles) to support supervised, auditable and accountable AI deployment. Rather than granting AI systems legal personhood, the strategy relies on human-centred legal frameworks and aims to anchor responsibility with institutions or individuals when AI is used.

Estonia's Information Systems Authority issued a 2024 report titled *Risks and Controls for Artificial Intelligence and Machine Learning Systems* (RIA 2024),⁷⁹ outlining structured methodologies for identifying cybersecurity, legal and societal risks. It provides concrete guidance on system description, threat modelling, applicable laws (especially GDPR and the EU AI Act) and control selection to mitigate harms (PrivacyCraft 2024).

According to the European Commission's AI Watch Estonia strategy report (European Commission 2019),⁸⁰ the country's executive plan includes the development of regulatory sandboxes for public-sector AI experimentation, reforms related to data governance and procurement, and institutional capacity-building measures such as establishing Chief Data Officers within ministries. These initiatives are intended to strengthen Estonia's ability to manage, oversee and coordinate AI deployment across government.

Foresight: regulation in Estonia / the European Union

Regulatory direction signals

- The EU AI Act (enforced from 2025-2026) is the strongest formal signal: a risk-based, rights-oriented model where agentic or autonomous AI would usually fall into high-risk categories when used in legal or administrative domains.
- The GDPR-AI Act interaction provides an established foundation for AI accountability, with well-defined rights to explanation, human oversight and redress.
- Estonia's own legal and digital reforms show a mature regulatory implementation culture, not merely compliance.

Technological and policy signals

- Estonia's KrattAI initiative demonstrates how interoperable public-sector AI ecosystems can remain under human control.
- Regular publication of AI risk and control frameworks by state bodies indicates continuous institutional learning.
- Continued discussions around "Krat Law" and algorithmic liability show openness to updating legal personhood and responsibility doctrines without conferring rights on AI.

Weak signals for 2030-2035

- Cross-border EU-UK or EU-US AI safety alignment could lead to mutual recognition of audits and certifications.
- Estonia may emerge as a testbed for AI assurance technologies (for example, blockchain-based audit trails).

4. AI regulatory context in China

China regulates artificial intelligence through a multi-layered, comprehensive framework built across data protection, cybersecurity, algorithm governance and scientific-ethics rules. At the core are three foundational laws: the Data Security Law 2021 (Standing Committee of the National People's Congress 2021a), the Personal

Information Protection Law 2021 (Standing Committee of the National People's Congress 2021b) and the Cybersecurity Law 2017 (Standing Committee of the National People's Congress 2017). These establish requirements for data handling, network security and oversight of information systems.⁸¹

Building on these, the Cyberspace Administration of China (CAC) has issued four AI-specific regulations: the Administrative of Algorithm-Generated Recommendations for Internet Information Services (2022), the Provisions on the Administrative of Deep Synthesis of Internet Information Services (2023), the Interim Measures for the Administration of Generative AI Services (2023) and the Measures for Labelling AI-Generated or Composed Content (2025). Ethical oversight is provided through the Trial Measures for Scientific and Technological Ethics Review (2023), issued by the Ministry of Science and Technology to guide research and technological development.⁸²

Beyond these core regulations, AI deployments must also comply with sector-specific laws, such as tort, consumer-protection, antitrust and criminal law, as well as specialised rules in industries like e-commerce and healthcare (for example, the E-Commerce Law (2019) and the Guidelines for Registration Review of AI Medical Devices). Collectively, this creates a broad regulatory ecosystem in which general data and cybersecurity laws provide the foundation, CAC regulations govern specific AI technologies, and sector-specific rules apply wherever AI is integrated into regulated domains.⁸³

China is complementing laws with national standards: in April 2025 regulators released three new standards aimed at improving generative-AI security and governance, scheduled to take effect 1 November 2025 (White & Case, 2025).⁸⁴

China's state-issued AI ethics frameworks, such as the *Ethical Norms for the New Generation Artificial Intelligence*, formally affirm principles of human decision-making primacy, fairness, controllability and accountability throughout the AI life-cycle (National New Generation Artificial Intelligence Governance Specialist Committee, 2021).⁸⁵ However, these ethics instruments operate within a governance system in which AI development and deployment must also align with broader political and ideological priorities. This becomes particularly clear in binding regulation: the 2023 *Interim Measures for the Management of Generative Artificial Intelligence Services* require providers not only to ensure accuracy, privacy protection and non-discrimination, but also to ensure that AI outputs conform to "social morality," "public order," and Core Socialist Values, while avoiding content deemed harmful to national security or social stability (CAC, 2023).⁸⁶

In the introduction to their book on *The Digital Silk Road* (2022), Gordon and Nouwens (2022) describe China's digital-connectivity initiative that emerged under the banner of the Belt and Road Initiative (BRI). The Digital Silk Road extends the idea of connectivity and infrastructure from the physical realm (roads, ports, rail) into the digital realm (telecommunications, data centres, digital economy, smart cities). While nested within the BRI, it has its own dynamics, driven by China's technological ambition, private tech firms and the global push to shape digital networks and standards. It signals China's increasing role in global digital infrastructure and suggests a potential shift in the architecture of the global cyber-economy.⁸⁷

Foresight: regulation in China

Regulatory direction signals

- The combination of data security, cybersecurity and generative AI content laws (2017-2024) already forms a state-centric, security-focused framework for AI governance.
- These measures demonstrate a centralised model where oversight of AI systems is embedded in state control and ideological supervision.
- Government-issued ethical guidelines highlight the importance of human oversight and fairness, but in practice they follow political priorities and censorship rules.
- AI regulation is closely linked to national security systems, showing that control over AI is firmly built into how the state governs.
- The three new national standards for generative AI (coming into effect in November 2025) show that China is moving from high-level rules to much more detailed technical requirements, including clearer expectations for data quality, risk management and how AI models must be tested.
- China's Digital Silk Road strategy indicates that its AI regulatory model may increasingly be exported through digital-infrastructure partnerships, influencing how data governance and AI standards are adopted in participating countries.
- As Chinese firms build telecommunications networks, data centres and cloud infrastructure abroad, compliance with China's own cybersecurity and AI-governance rules may become embedded in overseas digital systems.

Cultural and societal signals

- Trust in AI outcomes depends more on visible human and institutional oversight than on individual rights protections
- The development of ethical AI codes guided by the state reflects a focus on collective welfare and social order rather than personal autonomy.
- China's growing role in digital infrastructure abroad reinforces domestic confidence that its governance model is technologically successful and internationally competitive.
- The state's narrative of "safe and controllable" technology is strengthened by the global expansion of Chinese platforms and standards, shaping public perceptions of AI reliability and national technological leadership.

Weak signals for 2030-2035

- The existing laws could extend to cover more autonomous or agentic AI systems using the same state-driven regulatory mechanisms.
- Technical assurance and oversight standards for transparency, explainability and human supervision may gradually evolve within the current framework.
- Continued institutional integration of AI oversight into political and security systems is likely to persist as technology becomes more capable.
- AI governance may become more outward-facing, with China promoting its technical standards and regulatory frameworks through Digital Silk Road partnerships.

- International collaborations under the Digital Silk Road could lead to shared technical testing, certification and governance systems aligned with China's model.
- Countries adopting Digital Silk Road infrastructure might also adopt elements of China's AI governance approach, potentially extending the influence of state-centric, security-led regulatory norms into the global South.

End Notes

- ¹ Chan, A., Salganik, R., Markelius, A., Pang, C., Rajkumar, N., Krasheninnikov, D., Langosco, L., He, Z., Duan, Y., Carroll, M. and Lin, M., 2023, June. Harms from increasingly agentic algorithmic systems. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency* (pp. 651-666).
- ² Sapkota, R., Roumeliotis, K. I. & Karkee, M. (2025) 'AI Agents vs. Agentic AI: A Conceptual Taxonomy, Applications and Challenges', *Information Fusion*. Available at: <https://arxiv.org/abs/2505.10468>.
- ³ Government Digital Service, AI Insights – Agentic AI (10 February 2025, updated 5 June 2025) <https://www.gov.uk/government/publications/ai-insights/ai-insights-agentic-ai-html>
- ⁴ Judiciary of England and Wales (2025) Artificial Intelligence (AI): Guidance for Judicial Office Holders. 14 April. Courts and Tribunals Judiciary. Available at: <https://www.judiciary.uk/wp-content/uploads/2025/04/Refreshed-AI-Guidance-published-version.pdf>
- ⁵ Casper, S., Bailey, L., Hunter, R., Ezell, C., Cabalé, E., Gerovitch, M., Slocum, S., Wei, K., Jurković, N., Khan, A., Christoffersen, P., Ozisik, A., Trivedi, R., Hadfield-Menell, D. & Kolt, N. (2025) The AI Agent Index. arXiv preprint arXiv:2502.01635. Available at: <https://arxiv.org/pdf/2502.01635> Index available at <https://aiagentindex.mit.edu>
- ⁶ <https://aiagentindex.mit.edu/lucy/>
- ⁷ <https://aiagentindex.mit.edu/the-ai-scientist/>
- ⁸ Surden, H. (2024) 'Artificial intelligence and the practice of law: Risks and responsibilities', *Yale Journal of Law and Technology*, 26(3), pp. 101-138
- ⁹ North Carolina Central University (n.d.) Course Catalog: [Course details page]. NCCU eCatalog. Available at: https://ecatalog.nccu.edu/preview_course_nopop.php?catoid=26&coid=44237
- ¹⁰ Emory University School of Law (n.d.) AI and the Future of Work. Emory University. Available at: <https://law.emory.edu/centers-and-programs/ai-future-of-work.html>
- ¹¹ Salinski, D. (2025) 'In America: How Mobley v Workday Inc. Is Quietly Reshaping Legal and Legislative AI Accountability in the United States', *Modern Treatise*, 29 June. Available at: <https://www.moderntreatise.com/the-americas/2025/6/29/in-america-how-mobley-v-workday-inc-is-quietly-reshaping-legal-and-legislative-ai-accountability-in-the-united-states>
- ¹² U.S. District Court, Middle District of Florida, Case No. 433581.1.0, gov.uscourts.flmd.433581.1.0.pdf, The TMCA, February 2025.
- ¹³ Social Media Victims Law Centre (2025) AI Chatbot Lawsuit Clears Key Legal Hurdle in Landmark Case, Social Media Victims Law Centre: Blog, 23 May. Available at: <https://socialmediavictims.org/blog/character-ai-lawsuit-moves-forward>
- ¹⁴ Social Media Victims Law Centre (2025) AI Chatbot Lawsuit Clears Key Legal Hurdle in Landmark Case, Social Media Victims Law Centre: Blog, 23 May. Available at: <https://socialmediavictims.org/blog/character-ai-lawsuit-moves-forward>
- ¹⁵ Jasnow, D. (2025) 'New Lawsuits Targeting Personalized AI Chatbots Highlight Need for AI Quality Assurance and Safety Standards', *National Law Review*, 6 January. Available at: <https://natlawreview.com/article/new-lawsuits-targeting-personalized-ai-chatbots-highlight-need-ai-quality-assurance-and-safety-standards>
- ¹⁶ Jasnow, D. (2025) 'New Lawsuits Targeting Personalized AI Chatbots Highlight Need for AI Quality Assurance and Safety Standards', *National Law Review*, 6 January. Available at: <https://natlawreview.com/article/new-lawsuits-targeting-personalized-ai-chatbots-highlight-need-ai-quality-assurance-and-safety-standards>
- ¹⁷ The Law Society (2025) *Mazur v Charles Russell Speechlys: what it means for litigators*. Available at: <https://www.lawsociety.org.uk/topics/civil-litigation/mazur-v-charles-russell-speechlys-litigators>

-
- ¹⁸ Legal Services Board (2025) *Statement on High Court decision in Mazur v Charles Russell Speechlys LLP*. Available at: <https://legalservicesboard.org.uk/news/lsb-statement-high-court-decision-in-mazur-13-october-2025>
- ¹⁹ United States District Court for the Southern District of New York, *The New York Times Company v. Microsoft Corporation et al.*, No. 1:2023cv11195, Order Denying Motion for Reconsideration, May 16, 2025. Available at: <https://docs.justia.com/cases/federal/district-courts/new-york/nysdce/1:2023cv11195/612697/559>
- ²⁰ 'International Principles on Conduct for the Legal Profession' (International Bar Association, adopted on 28 May 2011, and updated on 11 October 2018. The Explanatory Note at para 9.2 in the Principles was further updated in May 2024) <https://www.ibanet.org/document?id=IBA-Legal-Technological-Competence-Statement-2024>
- ²¹ 21 O'Grady, C.G. and O'Grady, C., 2025. Agentic Workflows in the Practice of Law—AI Agents as Ethics Counsel. *Arizona Legal Studies Discussion Paper No. 25-03*; 38 *Georgetown Journal of Legal Ethics*, pp.247-. Available at: <https://ssrn.com/abstract=5113810>
- ²² Harasta, J., Novotná, T. & Savelka, J., 2024. It Cannot Be Right If It Was Written by AI: On Lawyers' Preferences of Documents Perceived as Authored by an LLM vs a Human. *arXiv preprint arXiv:2407.06798*. [online] Available at: <https://arxiv.org/abs/2407.06798>
- ²³ Solicitors Regulation Authority, 2024. Research shows trust in legal services high: Strategy benchmarking survey 2024. Solicitors Regulation Authority [online], 23 October. Available at: <https://www.sra.org.uk/sra/news/press/2024-press-releases/strategy-benchmarking-survey-2024/>
- ²⁴ The Law Society Gazette, 2024. Satisfaction levels and trust in lawyers continue to rise. *The Law Society Gazette*, [online] Available at: <https://www.lawgazette.co.uk/news/satisfaction-levels-and-trust-in-lawyers-continue-to-rise/5120359.article>
- ²⁵ Solicitors Regulation Authority, 2024. First Tier Complaints Report 2023. Solicitors Regulation Authority, 23 May [online] Available at: <https://www.sra.org.uk/sra/research-publications/first-tier-complaints-2023/>
- ²⁶ Fuller, S., Woodhall, C. & Walker, I., 2024. Making an impact. International Bar Association, 31 July [online], Available at: <https://www.ibanet.org/making-an-impact>
- ²⁷ Mukherjee, A. & Chang, H. H., 2025. Agentic AI: Autonomy, Accountability, and the Algorithmic Society. *arXiv preprint arXiv:2502.00289* [online] Available at: <https://arxiv.org/abs/2502.00289>
- ²⁸ Mansi, G. and Riedl, M. (2024) 'Recognizing lawyers as AI creators and intermediaries in contestability', *arXiv [Preprint]*, *arXiv:2409.17626*. doi: 10.48550/arXiv.2409.17626.
- ²⁹ Mukherjee, A. & Chang, T. (2025) Agentic AI: Autonomy, Accountability, and the Algorithmic Society. *arXiv preprint arXiv:2502.00289*. Available at: <https://arxiv.org/abs/2502.00289>
- ³⁰ UK Government (2025) The Bletchley Declaration by countries attending the AI Safety Summit, 1–2 November 2023. GOV.UK. Available at: <https://www.gov.uk/government/publications/ai-safety-summit-2023-the-bletchley-declaration/the-bletchley-declaration-by-countries-attending-the-ai-safety-summit-1-2-november-2023>
- ³¹ UK Government (2023) Safety Testing: Chair's Statement of Session Outcomes, 2 November 2023. GOV.UK. Available at: <https://assets.publishing.service.gov.uk/media/6544ec4259b9f5001385a220/aiss-statement-on-safety-testing-outcomes.pdf>
- ³² <https://assets.publishing.service.gov.uk/media/65c1e399c43191000d1a45f4/a-pro-innovation-approach-to-ai-regulation-amended-governement-response-web-ready.pdf>
- ³³ Department for Science, Innovation and Technology & AI Safety Institute (2025) International AI Safety Report 2025. GOV.UK. Available at: <https://www.gov.uk/government/publications/international-ai-safety-report-2025>
- ³⁴ UK Parliament (2025) Artificial Intelligence (Regulation) Bill [HL]. Available at: <https://bills.parliament.uk/bills/3942>

-
- ³⁵ <https://www.judiciary.uk/guidance-and-resources/artificial-intelligence-ai-judicial-guidance-2/>
- ³⁶ Crown Prosecution Service (2025) How the CPS will use Artificial Intelligence. GOV.UK. Available at: <https://www.cps.gov.uk/how-cps-will-use-artificial-intelligence>
- ³⁷ National Police Chiefs' Council (2023) Covenant for using Artificial Intelligence in Policing. Police Science and Technology. Available at: <https://science.police.uk/delivery/resources/covenant-for-using-artificial-intelligence-ai-in-policing/>
- ³⁸ College of Policing (2025) Responsible AI Checklist for Policing. College of Policing. Available at: <https://library.college.police.uk/docs/NPCC/Responsible-AI-checklist-for-policing-2025.pdf>
- ³⁹ The Law Society (2024) Generative AI: the essentials. [online] Available at: <https://www.lawsociety.org.uk/topics/ai-and-lawtech/generative-ai-the-essentials>
- ⁴⁰ Bar Council, Considerations when using ChatGPT and generative artificial intelligence software based on large language models (Information Technology Panel, 30 January 2024) <https://www.barcouncilethics.co.uk/documents/considerations-when-using-chatgpt-and-generative-ai-software-based-on-large-language-models>
- ⁴¹ Solicitors Regulation Authority (2019) Code of Conduct for Solicitors, RELs and RFLs. SRA. Available at: <https://www.sra.org.uk/solicitors/standards-regulations/code-conduct-solicitors/> (Accessed: 1 October 2025)
- ⁴² Solicitors Regulation Authority (2022) Compliance tips for solicitors regarding the use of AI and technology. SRA. Available at: <https://www.sra.org.uk/solicitors/resources/innovate/compliance-tips-for-solicitors/>
- ⁴³ Solicitors Regulation Authority (2023) Risk Outlook: The use of artificial intelligence in the legal market. SRA.
- ⁴⁴ Information Commissioner's Office (2023) Guidance on AI and Data Protection. ICO. Available at: <https://ico.org.uk/for-organisations/uk-gdpr-guidance-and-resources/artificial-intelligence/guidance-on-ai-and-data-protection/>
- ⁴⁵ Information Commissioner's Office (2020) Explaining decisions made with AI: Guidance for organisations. ICO. Available at: <https://ico.org.uk/for-organisations/uk-gdpr-guidance-and-resources/artificial-intelligence/explaining-decisions-made-with-artificial-intelligence/>
- ⁴⁶ Competition and Markets Authority (2024) AI Foundation Models: Update paper. GOV.UK. Available at: <https://www.gov.uk/cma-cases/ai-foundation-models-initial-review>
- ⁴⁷ Digital Regulation Cooperation Forum (2024) AI and Digital Hub - DRCF. Available at: <https://www.drcf.org.uk/ai-and-digital-hub>
- ⁴⁸ Digital Regulation Cooperation Forum (2024) AI and Digital Hub - DRCF [pilot]. Available at: <https://www.drcf.org.uk/ai-and-digital-hub>
- ⁴⁹ Scottish Government. (2021) *Scotland's artificial intelligence strategy: Trustworthy, ethical and inclusive*. gov.scot. Available at: <https://www.gov.scot/publications/scotlands-ai-strategy-trustworthy-ethical-inclusive/>
- ⁵⁰ Northern Ireland Assembly Research and Information Service. (2024) *Artificial Intelligence regulation, use and innovation in the United Kingdom: a broad overview* (Paper 11/24, NIAR 45-24). Northern Ireland Assembly. Available at: <https://www.niassembly.gov.uk/globalassets/documents/raise/publications/2022-2027/2024/economy/1124.pdf>
- ⁵¹ Northern Ireland Executive. (2025) "New AI and Digital Office will drive innovation and transform services - O'Neill and Little-Pengelly." Available at: <https://www.executiveoffice-ni.gov.uk/news/new-ai-and-digital-office-will-drive-innovation-and-transform-services-oneill-and-little-pengelly>
- ⁵² Loring, J.M., 2025. When AI Acts Independently: Legal Considerations for Agentic AI Systems. *The National Law Review*, 10 June [online] Available at:

<https://www.natlawreview.com/article/when-ai-acts-independently-legal-considerations-agentic-ai-systems>

⁵³ KPMG (2025) State series: AI legislation regulatory alert. KPMG. Available at: <https://kpmg.com/kpmg-us/content/dam/kpmg/pdf/2025/state-series-ai-legislation-reg-alert.pdf>. MultiState (2025) AI legislative updates, vol. 71. MultiState. Available at: <https://www.multistate.ai/updates/vol-71>

⁵⁴ MultiState (2025) AI legislative updates, vol. 71. MultiState. Available at: <https://www.multistate.ai/updates/vol-71>

⁵⁵ TrustArc (2025) Utah's AI Policy Act Is Here. Is Your AI Ready to Confess? TrustArc. Available at: <https://trustarc.com/resource/utah-ai-policy-act>

⁵⁶ Utah State Legislature (2024) Senate Bill 149: Generative Artificial Intelligence Disclosure Requirements, Chapter 2, Section 12. Available at: <https://le.utah.gov/~2024/bills/static/SB0149.html#13-2-12>

⁵⁷ Utah Legislature (2024) SB 149 Artificial Intelligence Policy Act. Utah State Legislature. Available at: <https://le.utah.gov/~2024/bills/sbillint/SB0149.htm>

⁵⁸ BCLP (2024) California turns to the use of AI in healthcare. Bryan Cave Leighton Paisner LLP. Available at: <https://www.bclplaw.com/en-US/events-insights-news/california-turns-to-the-use-of-ai-in-healthcare.html>

⁵⁹ Seyfarth (2024) Artificial intelligence legal roundup: Colorado postpones implementation of AI law. Seyfarth Shaw LLP. Available at: <https://www.seyfarth.com/news-insights/artificial-intelligence-legal-roundup-colorado-postpones-implementation-of-ai-law-as-california-finalizes-new-employment-discrimination-regulations-and-illinois-disclosure-law-set-to-take-effect.html> (Accessed: 1 October 2025).

⁶⁰ Brookings (2025) How different states are approaching AI. Brookings. Available at: <https://www.brookings.edu/articles/how-different-states-are-approaching-ai/>

⁶¹ Examples of uses in legislation from California, Colorado, Kentucky, Maryland and Utah have been summarised in the previous section.

⁶² https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=202320240AB2885

⁶³ Anderson, H., Reem, N. & Susas, J. (2025) Automated Decision Making Emerges as an Early Target of State AI Regulation. White & Case LLP (Insight Alert), 7 March. Available at: <https://www.whitecase.com/insight-alert/automated-decision-making-emerges-early-target-state-ai-regulation>

⁶⁴ Anderson, H., Reem, N. & Susas, J. (2025) Automated Decision Making Emerges as an Early Target of State AI Regulation. White & Case LLP (Insight Alert), 7 March. Available at: <https://www.whitecase.com/insight-alert/automated-decision-making-emerges-early-target-state-ai-regulation>

⁶⁵ Executive Office of the President, 2023. Executive Order 14110: Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. Federal Register, 1 November 2023, 88 FR 75191–75226 [online] Available at: <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>

⁶⁶ Trump, D.J., 2025. Removing Barriers to American Leadership in Artificial Intelligence (Executive Order 14179). Federal Register, 31 January, 90 FR 8741–8742 [online] Available at: <https://www.federalregister.gov/documents/2025/01/31/2025-02172/removing-barriers-to-american-leadership-in-artificial-intelligence>

⁶⁷ Federal Trade Commission (2025) 'AI and the Risk of Consumer Harm', Tech at FTC, 3 January. https://data.aclum.org/storage/2025/01/FTC_www_ftc_gov_policy_advocacy-research_tech-at-ftc_2025_01_ai-risk-consumer-harm.pdf

-
- ⁶⁸ Tabassi, E. (2023) Artificial Intelligence Risk Management Framework (AI RMF 1.0), National Institute of Standards and Technology, Gaithersburg, MD. Available at: <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>
- ⁶⁹ Geistfeld, M. (2024) Principles of the Law, Civil Liability for Artificial Intelligence, The American Law Institute. Available at: <https://www.ali.org/project/principles-law-civil-liability-artificial-intelligence>
- ⁷⁰ American Bar Association (2024) Formal Opinion 512: Lawyers' Use of Generative Artificial Intelligence Tools. ABA Standing Committee on Ethics and Professional Responsibility. Available at: <https://www.lawnext.com/wp-content/uploads/2024/07/aba-formal-opinion-512.pdf>
- ⁷¹ Hepburn, J. D. (2024) 'Florida Bar Provides Guidance on Lawyers' Use of Generative AI', McLane Middleton Insights, 20 March. Available at: <https://www.mclane.com/insights/florida-bar-provides-guidance-on-lawyers-use-of-generative-ai/>
- ⁷² European Commission, AI Watch (2025) *Estonia AI strategy report*. Available at: https://ai-watch.ec.europa.eu/countries/estonia/estonia-ai-strategy-report_en
- ⁷³ e-Residency of Estonia (2025) *e-Residency of Estonia - home page*. Available at: <https://www.e-resident.gov.ee>
- ⁷⁴ Kaspar Korjus (2018) 'Estonian president Kersti Kaljulaid reveals the future direction of e-Residency', *e-Residency blog*, 18 December. Available at: <https://www.e-resident.gov.ee/blog/posts/estonian-president-kersti-kaljulaid-reveals-the-future-direction-of-e-residency-2/>
- ⁷⁵ Estonian Information System Authority (RIA) (2025) *Bürokratt: interoperable network of public-sector virtual assistants*. Available at: <https://www.ria.ee/en/state-information-system/personal-services/burokratt>
- ⁷⁶ e-Estonia (2021) *National AI strategy–KrattAI factsheet*. Available at: <https://e-estonia.com/new-e-estonia-factsheet-national-ai-kratt-strategy/>
- ⁷⁷ European Commission (2023) *Digital public services based on open source: Case study–Bürokratt*. Available at: <https://interoperable-europe.ec.europa.eu/collection/open-source-observatory-osor/document/digital-public-services-based-open-source-case-study-burokratt>
- ⁷⁸ Ministry of Economic Affairs and Communications, Estonia (2021) *Kratt strategy: Estonia's national AI strategy 2022-2023*. Available at: https://e964029b-7c32-42a7-817c-37c391a3b976.filesusr.com/ugd/980182_4434a890f1e64c66b1190b0bd2665dc2.pdf
- ⁷⁹ Estonian Information System Authority (RIA) (2024) *Risks and controls for artificial intelligence and machine learning systems*. Tallinn: RIA. Available at: <https://www.ria.ee/sites/default/files/documents/2024-05/Risks-and-controls-for-artificial-intelligence-and-machine-learning-systems.pdf>
- ⁸⁰ European Commission AI Watch (2019) *Estonia AI strategy report*. Available at: https://ai-watch.ec.europa.eu/countries/estonia/estonia-ai-strategy-report_en
- ⁸¹ Chambers & Partners (2025) 'Artificial intelligence 2025 - China'. *Global Practice Guide*. Available at: <https://practiceguides.chambers.com/practice-guides/artificial-intelligence-2025/china>
- ⁸² Chambers & Partners (2025) 'Artificial intelligence 2025 - China'. *Global Practice Guide*. Available at: <https://practiceguides.chambers.com/practice-guides/artificial-intelligence-2025/china>
- ⁸³ Chambers & Partners (2025) 'Artificial intelligence 2025 - China'. *Global Practice Guide*. Available at: <https://practiceguides.chambers.com/practice-guides/artificial-intelligence-2025/china>
- ⁸⁴ White & Case LLP 2025. *AI Watch: Global Regulatory Tracker - China*. Available at: <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-china>

⁸⁵ National New Generation Artificial Intelligence Governance Specialist Committee. 2021. *Ethical Norms for New Generation Artificial Intelligence* (trans. Centre for Security and Emerging Technology). Available at: CSET, Georgetown University (pdf) https://cset.georgetown.edu/wp-content/uploads/t0400_AI_ethical_norms_EN.pdf

⁸⁶ CAC (Cyberspace Administration of China) 2023. *Interim Measures for the Management of Generative Artificial Intelligence Services*. Available at: <https://www.chinalawtranslate.com/en/generative-ai-interim/>

⁸⁷ Gordon, D. & Nouwens, M. (2022) 'Introduction The Digital Silk Road: China's technological rise and the geopolitics of cyberspace', *International Institute for Strategic Studies (IISS), Adelphi Online Analysis*, 6 December. Available at: <https://www.iiss.org/online-analysis/online-analysis/2022/12/digital-silk-road-introduction>