

PREPARING FOR AI:

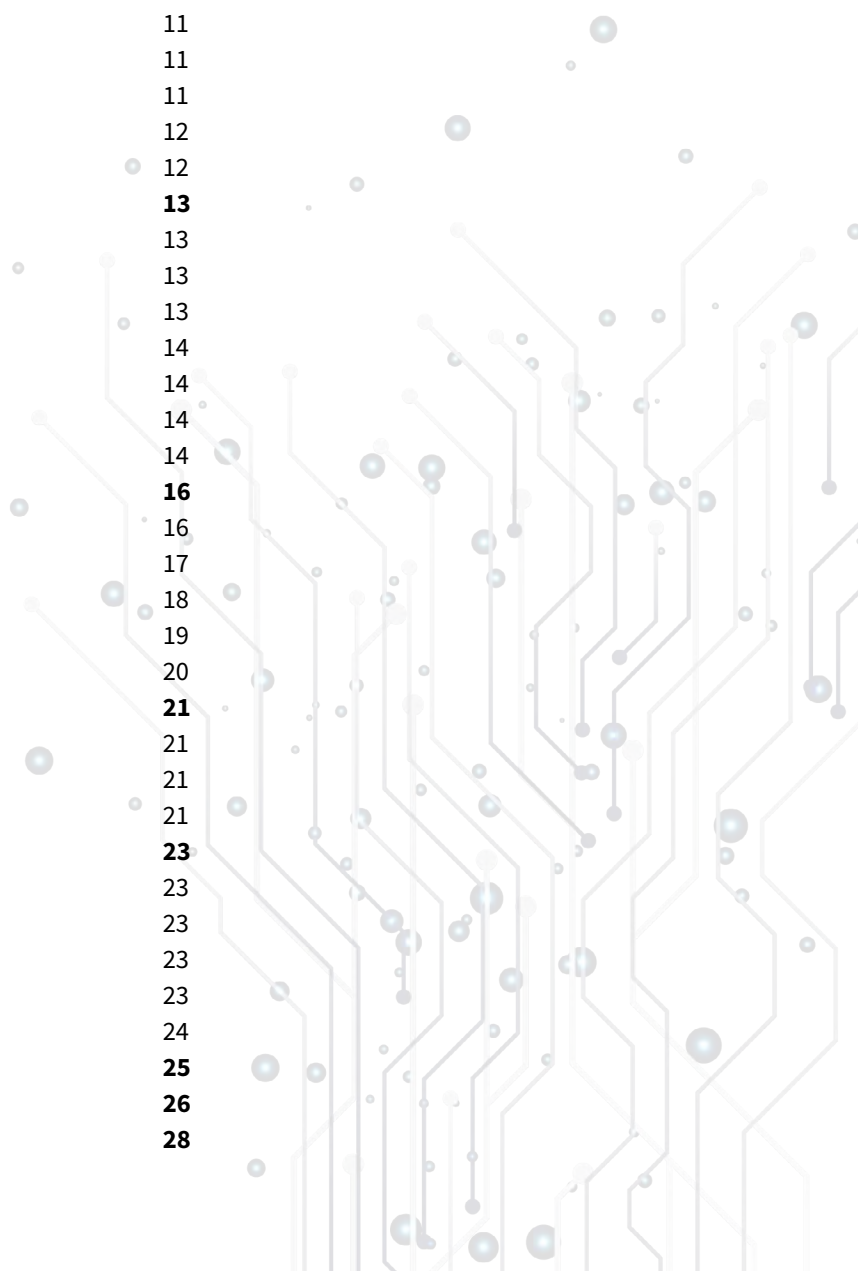
A Readiness Blueprint
for Regulators

A white paper for regulators, informed by regulators.

GRCS PARTNERS INC.

Contents

Executive Summary	3
Three actions to start this quarter	3
Contributors	4
Beyond the blueprint	4
1. Introduction: Why AI readiness matters for regulators	5
1.1 The shift regulators cannot ignore	5
1.2 The regulator's context	6
1.3 Purpose and scope of this blueprint	6
1.4 From blueprint to practice: What success looks like	7
2. Understanding the AI landscape	8
2.1 Plain-English primer	8
2.2 Common regulatory use cases	8
2.3 Shadow AI and embedded AI	9
2.4 Principles over perfection	9
Concluding note on the landscape	10
3. Core pillars of AI readiness	11
3.1 Governance	11
3.2 Policy and compliance	11
3.3 People and culture	11
3.4 Technology and controls	12
From pillars to practice	12
4. Applications in regulatory practice	13
4.1 Licensing and registration	13
4.2 Complaints and enquiries	13
4.3 Inspections and fieldwork	13
4.4 Investigations and enforcement	14
4.5 Policy and guidance	14
4.6 High-risk domains: healthcare and justice	14
From possibilities to practice	14
5. Practical pathways to get started	16
5.1 The 30–60–90 plan	16
5.2 Pilot design and measurement	17
5.3 Capability uplift program	18
5.4 Procurement fast lanes	19
On practical pathways	20
6. Consequences of inaction	21
6.1 Risk taxonomy	21
6.2 Scenario snapshots	21
6.3 Public trust and communications	21
7. Recommendations and next steps	23
7.1 Immediate actions–next 30 days	23
7.2 Next 100 days	23
7.3 One-year roadmap	23
7.4 Collaborate across jurisdictions	23
Recommendations: closing note	24
Conclusion–and looking ahead	25
Contributors	26
Authors	28



Executive Summary

Artificial intelligence has moved from experimentation to enterprise utility. For regulators, the question is not whether to use AI, but how to use it safely, visibly, and to deliver clear public value. Recent studies show both momentum and a maturity gap: 64% of public-sector organizations are already exploring or actively working on generative-AI initiatives, yet only 21% have moved to pilots or deployment; fewer than one-quarter report high maturity in any aspect of data readiness, even as 71% cite efficiency and cost savings as key incentives. Looking ahead, 90% plan to explore, pilot, or implement agentic AI within the next two to three years—underscoring the need to build guardrails and data foundations now, not later.¹

This blueprint proposes a measured path forward. It sets out four pillars of readiness—governance, policy and compliance, people and culture, and technology and controls—plus concrete sector examples, a 30–60–90 plan to get started, and checklists to make decisions repeatable. Principle-based oversight, continuous assurance, and visible “human-in-the-loop” decision boundaries anchor the approach. The goal is to minimize low-value toil and accelerate quality, while maintaining trust.

Leaders interviewed for this paper highlighted both urgency and pragmatism. “If you’re not doing it now, then you’re late to the game,” said Norton. Paul Byrne reminded us that “trust is a regulator’s currency, so you have to spend wisely.” Joe McIntosh warned that “we’re using a 1990s approach to solve a 2030s problem,” urging near-real-time monitoring over periodic reviews. Several emphasized scale and inevitability: “Assume up to half your people are already using AI in some capacity,” said George Plytas, and Steve Genders observed that AI is “democratizing intelligence... shortening those decision flows.”

A measured posture—neither hype-driven nor paralyzed—wins. UK analysis suggests that around 41% of public-sector work time is exposed to generative AI (about 3.4 hours in an 8.3-hour day), yet only 5% of activities are fully replaceable.² This frames the opportunity clearly: AI is an accelerant, not a substitute. And tangible movement is already under way—Canada has identified around 300 federal AI projects, most within agencies, showing that the shift is real.

As Byrne put it, “from a risk-based approach, the risk of not taking this action is higher than the risk of being ultra cautious.” This paper offers a blueprint to act now, learn fast, and show your work.

Success over the next 12 months looks like: fewer low-value manual steps in routine work; explicit guardrails staff actually use; transparent notes on where AI assists and where humans decide; and a repeatable cadence of review and improvement. This paper shows how to get there, starting small, proving value, and scaling safely.

Three actions to start this quarter

1. Form a cross-functional AI governance group with a simple charter, a monthly cadence, and an AI-use inventory as its first artefact.
2. Publish interim guardrails (acceptable use, prohibited use, approvals, redaction, transparency), and schedule a one-hour introductory course for staff.
3. Select a narrow, internal pilot with clean data, clear exit criteria, and visible human-in-the-loop controls; measure before and after.

¹ Capgemini Research Institute, *Data Foundations for Government: From AI Ambition to Action*, 2025.

² Alan Turing Institute, *Mapping the Potential: Generative AI and Public Sector Work*, 2025.

Contributors

Artificial intelligence has vastly expanded access to knowledge, but the quality of that knowledge can vary. Open sources, reports, and commentary are plentiful, yet separating durable insights from hype or misconception requires careful work. To ground this paper in what is both true and important for regulators, we spent many hours in conversation with experts and leaders at the intersection of AI and regulation. Notably, Byrne—whose risk-based perspective on governance is cited throughout—along with practitioners such as Wade Hillier, Joe McIntosh, George Plytas, Bernie Plourde, Heather Kelley, Andrew Norton, Ronne Hines, and Steve Genders. Their experience, combined with public-sector research on adoption and readiness, shaped the blueprint you see here.

Interviews

Nine leaders across regulation, technology, and information security were interviewed using a semi-structured approach: Paul Byrne, Wade Hillier, Joe McIntosh, George Plytas, Bernie Plourde, Heather Kelley, Andrew Norton, Ronne Hines, and Steve Genders. Quotations included here are verbatim excerpts used as illustrative evidence.

Research

We reviewed current public-sector studies on AI adoption and risk, focusing on adoption gaps, training and policy awareness, data and infrastructure constraints, and emerging regulatory frameworks (EU AI Act, Canada's federal directives and strategy, Ontario's directive, Australia's lifecycle standard).

Triangulation

Themes were cross-checked across interviews and research. Where multiple sources agreed (e.g., the importance of governance cadence, data readiness, and visible human checks), those points are emphasized.

Contributors



Paul Byrne

Regulator; risk-based governance and human-in-the-loop accountability.



Wade Hillier

Regulator/executive advisor; data-driven decision-making and pragmatic guardrails.



Joe McIntosh

Technology leader; continuous assurance and co-pilots for oversight.



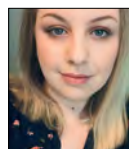
George Plytas

Operations/IT leader; governance-first, ISO/NIST alignment, platform standardization.



Bernie Plourde

Policy leader; staff guardrails and administrative automation boundaries.



Heather Kelley

Licensing operations; intake automation, validation, and training cadence.



Andrew Norton

CIO; pilot gates, tenant-level security, shadow-AI clamp.



Ronne Hines

Agency leadership; incremental pilots, privacy inventory.



Steve Genders

Information security; endpoint-first controls, observability, FAIR risk quantification.

Beyond the blueprint

Each of the pillars and pathways introduced in this paper—governance, continuous assurance, procurement, sector-specific applications, and measurement—merits its own in-depth treatment. GRCS will be publishing a series of follow-on papers that take these topics further, offering deeper dives, caselets, and practical tools for regulators who want to move from blueprint to detailed playbook. What follows here is a high-level map to readiness, designed to help agencies take first steps with confidence while knowing more detailed guidance is on the way.

1. Introduction: Why AI readiness matters for regulators



“If a humanoid robot drops a patient in a care home, is it an industrial accident? Is it the operator? Is it the software? Is it the designer of the system?”

– Paul Byrne, on the complexity of AI accountability in healthcare

It’s a vivid example, and one that grounds often abstract conversations about AI firmly in reality. Accountability, as Byrne suggests in this example, is hard to pin down. It doesn’t live in a single place and instead straddles the operator who clicked “accept,” the system that generated the advice, the data that fed it, and the institution that chose to use it.

That is the regulator’s world now: decisions increasingly arrive with an invisible assist, and the chain of responsibility must still be visible, auditable, and fair. Byrne suggested that the risk of not taking this action is greater than the risk of abundant caution. If trust, as he likes to say, is “a regulator’s currency,” it would surely deplete rapidly if an agency couldn’t demonstrate that people—not machines—are in control at all times.

Our interviews with front-line leaders echoed two themes: urgency and inevitability. Norton’s view was unvarnished: regulators that defer action will be playing catch-up. Hillier cautioned against denial: “You can’t call yourself a modern regulator if you’re saying you can’t use AI, because that seems crazy.” Even without formal rollouts or compelling, large-scale examples of AI in action in a regulatory setting, AI has already slipped into day-to-day work: staff draft, summarize, and triage with help from tools embedded in the platforms they already use; licensees are doing the same. As McIntosh noted, AI has crossed into everyday tools used by staff and the public alike.

Urgency, however, is not a license to rush. The leaders we spoke with argued for a measured posture that is neither hype-driven nor paralyzed. Measured means clear roles and decision rights, visible gates and approvals, unambiguous boundaries where humans decide, and transparent notes that show how AI assists rather than replaces professional judgment. It also means being realistic about capacity. As Wade

Hillier noted, smaller regulators “can adapt faster, but bandwidth is tight”—all the more reason to start with narrow pilots that deliver obvious value and lower risk. Kelley’s practical lens—reduce low-value toil at the front end so experts can focus on timeliness and quality—captures the spirit of the work.

This paper is built from those conversations. It distills practitioner insight into a readiness blueprint: principle-based guardrails you can adopt now, a 30–60–90 plan to get started, and sector-grounded examples of where AI can assist while people remain accountable. The aim is simple: help regulators move early, learn fast, and show their work—so the public and licensees can see that trust is being earned, not assumed.

1.1 The shift regulators cannot ignore

AI has moved rapidly from background experiments to consumer-grade assistance, embedded in mainstream tools and platforms. As McIntosh observed, it has “moved from the background to the forefront and into consumers’ hands.” For regulators with mandates rooted in protection and trust, that shift matters. Even if adoption in regulatory agencies remains limited, the reality is that staff in your organizations—and the licensees you oversee—are already beginning to experiment with AI to draft, summarize, triage, and decide.

Concrete examples are rare, but they are starting to appear—if targeted in scope. In Ohio, regulators used AI to comb through the state’s Administrative Code, eliminating over two million words of outdated or redundant content—work that would have taken human teams years to complete.³ Employment and Social Development Canada has applied AI to automatically

³ Sarah Donaldson, *Artificial intelligence axing outdated terms in Ohio regulatory code*, The Statehouse News Bureau, 27 Dec 2024.

check Record of Employment forms for completeness and errors before human review.⁴ And at the U.S. Internal Revenue Service, AI models are already helping to flag anomalies in tax returns, triaging high-risk cases for auditors.⁵ Each of these is narrow in scope, but together they illustrate how AI can begin to shift regulatory practice.

The traditional regulatory rhythm of periodic checks followed by investigation is unlikely to keep pace with AI's speed and variability. McIntosh put the contrast bluntly: "Regular check-ins aren't the same as near-real-time analysis and tools telling you where there's a problem." In short, standing still is not a neutral choice; it heightens both operational and reputational exposure.

Evidence beyond regulation also points to the scale of the shift. A recent UK government trial involving more than 20,000 civil servants found that AI tools such as Microsoft Copilot saved staff an average of 26 minutes per working day—about two weeks a year—on administrative tasks like drafting and meeting summaries, and that over 80% of participants wanted to continue using the technology.⁶ As outlined in the introduction, UK analysis indicates that a large share of public-sector work now falls within generative AI's reach, but only a small subset of tasks are suitable for full automation. The message is clear: AI is not a wholesale substitute but an accelerant, speeding up the work regulators already do.

Interviewees for this paper stressed that governing such a fast-moving technology will be difficult. "It's going to be super hard: it's so dynamic, so agile that figuring out how you can govern it is very hard," cautioned Genders. The response, argued Norton, is not to retreat but to proceed carefully: "It's not cautious, it's measured." A measured approach means setting clear roles and decision rights, establishing visible gates, defining when humans must review and decide, and communicating transparently.

1.2 The regulator's context

Any discussion of AI readiness has to start with the realities regulators face. Agencies operate under heightened transparency expectations, often with constrained budgets and uneven levels of digital maturity across their own teams and among the licensees they oversee. Some regulators run modern,

cloud-based systems; others still rely on paper files and fax machines. That inconsistency shapes what regulators can reasonably attempt.

Several interviewees underscored just how manual many processes remain. Kelley described "a lot of paper, a lot of manual processing," pointing to opportunities for AI to streamline intake, triage, and drafting so professionals can focus on judgment, consistency, and timeliness rather than clerical work. Hillier noted that "smaller regulators can adapt faster, but bandwidth is tight," suggesting that early AI pilots should be scoped narrowly, with clear value and limited risk, to avoid overextending scarce resources.

Capacity constraints go beyond staff time. Many regulators lack the data foundations or technical platforms to deploy AI safely. Legacy systems are common, data is often siloed, and governance structures are still catching up to cloud adoption. At the same time, expectations from the public and government are rising. Citizens compare regulatory services to the best digital experiences they encounter elsewhere, and political leaders expect efficiency gains from AI even before regulators have the tools or staff to deliver them.

This mix of high expectations, limited resources, and uneven maturity is the context in which AI readiness must be planned. It explains both why adoption has been slow to date and why a measured, capability-focused approach—starting small, proving value, and scaling carefully—is the only viable path forward.

1.3 Purpose and scope of this blueprint

This blueprint is meant to be practical: a tool regulators can use to bring AI into everyday operations safely, transparently, and with public value at the center. It focuses on areas like productivity, records, licensing, inspections, and evidence handling, while steering clear of technical build recipes or vendor prescriptions. The aim is progress regulators can point to within a year, guided by adaptable standards and patterns rather than rigid rules.

Purpose

Offer regulators practical guardrails and accelerators to use AI safely, transparently, and with public value in mind.

⁴ Government of Canada, *Responsible Use of Artificial Intelligence in the Government of Canada* (AI use case inventory), 2025.

⁵ FedScoop, *IRS turning to AI for fraud detection and enforcement*, 2024.

⁶ Financial Times, *UK civil service trial finds AI tools save staff two weeks a year*, 2024.

Scope

Everyday regulatory work such as internal productivity, records and correspondence, licensing and inspections preparation, evidence and complaints intake, and analysis support.

Out of scope

Detailed technical build recipes or vendor-specific solutions; this paper points instead to adaptable standards and patterns.

1.4 From blueprint to practice: What success looks like

As a practical tool regulators can use, the real measure of success is whether it helps agencies make progress they can point to within a year.

What success looks like in 12 months:

- **Outcomes:** faster turnaround on routine tasks, guardrails that are clear and consistent, and visible human-in-the-loop decisions.
- **Indicators:** fewer low-value manual steps, fewer policy exceptions, and transparent reports on where and how AI is used.
- **Habits:** a monthly governance cadence, continuous-assurance signals that surface exceptions early, and role-based training supported by communities of practice.



2. Understanding the AI landscape

Artificial intelligence is not one technology but a family of capabilities moving at different speeds. For regulators, clarity matters. The conversation is often blurred by hype, buzzwords, and shifting definitions. As Byrne put it, “most current debates on AI are really about large language models. Precision in naming the systems under review is the foundation of governing them wisely.”

Regulators do not need to master every technical detail, but they do need a working map of the

landscape. They need to know what tools exist, how they are already being used, and what risks or opportunities each class of tool brings.

2.1 Plain-English primer

The primer that follows is intended to cut through jargon. It gives regulators and boards a way to talk about AI in plain terms, without confusion or unnecessary complexity.

AI Model	Definition	Example
Large Language Models (LLMs)	Systems that generate and transform text. They excel at drafting, summarizing, and classifying when given the right context.	ChatGPT (OpenAI), Claude (Anthropic), Gemini (Google DeepMind), DeepSeek, LLaMA (Meta), Mistral, Falcon (TII), Cohere Command R, AI21 Jamba, Grok (xAI)
Retrieval-augmented assistants	LLMs linked to an organization’s own documents for grounded answers, avoiding reliance on generic internet data.	Perplexity; enterprise tools such as Copilot, where connection to internal data would require careful consideration.
Classification models	Simpler models trained to tag or route cases, documents, or messages; useful for triage.	Spam filter, OCR tagger
Lightweight agents	Scripts that chain steps (retrieve → draft → check) with human approval, automating processes without removing oversight.	Zapier bot, Copilot Studio

The purpose of this primer is not to make regulators technologists, but to provide an easy reference for a shared language. Clarity reduces confusion, keeps expectations realistic, and makes governance discussions concrete.

2.2 Common regulatory use cases

With definitions clear, the next step is recognizing where AI is already useful (or could be useful) in regulatory settings. These use cases are not speculative; some are live today inside agencies, sometimes formally sanctioned, sometimes appearing as shadow AI.

- **Internal productivity:** drafting letters, summarizing case notes, preparing meeting minutes.
- **Records and correspondence:** sorting inboxes, standardizing subject lines, routing common enquiries.

- **Complaints trend analysis:** clustering themes, flagging surges, proposing FAQs.
- **Licensing workflow:** validating standard documents, generating requests for missing information, preparing batch approvals where criteria are met.
- **Inspections preparations:** pre-visit packs summarizing prior history, inspection checklists aligned to risk factors.
- **Evidence intake:** transcription, de-identification, and structuring of submissions.

In all these areas, AI *supports* rather than replaces professional judgment. As Byrne reminded us, “accountability remains with humans and not the algorithm...any decision that impacts humans is made by a human.”

2.3 Shadow AI and embedded AI

Shadow AI and embedded AI are already shaping how regulators and their staff encounter this technology day to day. Some uses emerge when individuals sign up for tools on their own, others arrive quietly inside the software platforms agencies already rely on. Both carry risks if left unmanaged, but they require different responses.

Shadow AI

The use of AI tools by staff without organizational approval or oversight. Examples include employees signing up for ChatGPT with personal accounts, or pasting agency documents into free online tools. It is “shadow” because it operates outside governance, security, and record-keeping processes.

Embedded AI

AI features built into software and platforms already in use, often switched on by default by vendors. Examples include Microsoft 365 Copilot suggesting text in Word, or Salesforce adding predictive fields. It is “embedded” because it arrives inside existing systems, sometimes without explicit procurement or review.

Staff sign-ups to unapproved tools and vendor-enabled features are already present in most environments. Norton described finding shadow subscriptions and “cutting them off and offering Copilot as the alternative.” Genders warned that blocking alone backfires: “There have been stories of people who’ve tried [to shut it off] and they go around. There’s so much incentive to have that intelligence boost.”

The lesson is clear: prohibitions without alternatives drive workarounds. A discovery-to-sanctioned-use pathway is essential: detect, evaluate, approve with guardrails, then monitor.

2.4 Principles over perfection

For regulators, the pace of AI development will always outstrip the pace of formal lawmaking. That makes principle-based guardrails essential—flexible enough to adapt, but firm enough to provide clarity and accountability.

Fortunately, regulators do not need to start from scratch. International standards already exist to guide responsible AI use:

- **ISO 42001:** a new international standard that sets out how organizations can build and run an AI management system, applying the discipline of ISO-style governance to the unique risks of AI.⁷
- **NIST AI Risk Management Framework (AI RMF):** a practical playbook from the US National Institute of Standards and Technology that helps organizations identify, assess, and mitigate AI risks across design, development, and deployment.⁸
- **NIST Cybersecurity Framework (CSF) 2.0:** a voluntary framework from the US National Institute of Standards and Technology that helps organizations strengthen cybersecurity, manage evolving risks, and improve resilience through clear, outcome-focused guidance.

Together, these frameworks offer regulators a head start and a set of common reference points to draw from. Because AI tools evolve faster than formal regulations, regulators need principle-based guardrails rather than rigid rulebooks. As Byrne argued, “I wouldn’t make it a policy; I’d make it a framework...align it to the AI ISO standard.” Plytas recommended, “cherry-pick from ISO 42001 and NIST AI RMF for your regulatory stack.”

Guardrails or leashes?

Cary Coglianese, a leading scholar of regulation, has suggested that “leashes” may be a more apt metaphor than “guardrails”—emphasizing that AI must remain firmly tethered to human control, not merely bounded by external limits.⁹ It’s a valuable reminder that metaphors shape how we think about governance. In this paper, we use guardrails in a pragmatic sense: clear, principle-based boundaries that help regulators and staff understand what is permitted, what is prohibited, and when human judgment must apply.

Internationally, the regulatory direction of travel is clear. The EU AI Act establishes a tiered risk model: minimal-risk uses with no restriction, higher-risk uses subject to transparency and monitoring, and high-risk public-sector applications (such as in justice or border control) requiring impact assessments, disclosure, and mandatory human oversight.¹⁰

⁷ International Organization for Standardization, *ISO/IEC 42001:2023—Artificial intelligence management system*, 2023.

⁸ National Institute of Standards and Technology, *AI Risk Management Framework (AI RMF 1.0)*, 2023.

⁹ Cary Coglianese, *Leashes, Not Guardrails: The Future of AI Regulation*, Risk Analysis, 2024.

¹⁰ European Parliament and Council of the European Union, *Regulation (EU) 2024/1689 laying down harmonized rules on artificial intelligence (AI Act)*, adopted 21 May 2024.

National exemplars are emerging:

- Canada’s 2025–27 federal AI strategy requires mandatory Algorithmic Impact Assessments and public registries of AI systems.¹¹
- Australia’s lifecycle technical standard and GovAI collaboration platform aim to ensure AI systems are monitored from design to decommission, with regulators and vendors sharing a common workspace.^{12 13}

These frameworks illustrate “what good looks like” and provide models regulators can adapt without waiting for perfection. The destination is continuous assurance—moving from periodic spot checks to signals and dashboards that surface exceptions early.

Concluding note on the landscape

Understanding the AI landscape is the foundation of readiness. Regulators do not need to become AI engineers, but they do need a shared language, a clear view of common use cases, and awareness of the risks of unmanaged adoption. The lesson from international exemplars is not that every jurisdiction must regulate identically, but that principles—risk tiers, transparency, human oversight, lifecycle monitoring—travel across borders. With clarity on what AI is and how it is being used, regulators can move from reactive bans and shadow workarounds to measured adoption, backed by frameworks that evolve as the technology does.



“History tells us that every new technology brings a defining incident. Trust won’t be lost by the first major AI event. It will be lost if regulators cannot show they were ready to respond with clarity and control.”

– Paul Byrne, on the inevitability of AI-related incidents

¹¹ Government of Canada, *Directive on Automated Decision-Making*, Treasury Board of Canada Secretariat, 2019 (last modified 2023).

¹² *Australian Government AI Technical Standard*, Digital Transformation Agency (DTA), 2025.

¹³ *GovAI platform launch for the Australian Public Service*, Government of Australia, 31 July 2025.

3. Core pillars of AI readiness

Getting AI readiness right is less about a single tool and more about building capacity across four connected pillars: governance, policy and compliance, people and culture, and technology and controls. Governance sets direction and keeps accountability clear. Policy and compliance turn principles into daily rules. People and culture determine whether adoption sticks. Technology and controls provide the guardrails for safe use.

These pillars are not abstract: They draw directly on regulator insights and international findings. The

EU AI Act, for example, embeds human-in-the-loop obligations for high-risk uses. Canada requires agencies to publish impact assessments and AI registries. Many public servants have had some training, yet many do not know whether their organization has a policy. Across the board, ambition is constrained by legacy data and infrastructure. The following sections describe each pillar according to interviews with our contributors and wider research.

Trusted AI in Regulation				
Pillars	Governance	Policy and Compliance	People and Culture	Technology and Controls
	• Roles and decision rights • Cadence and artefacts • Human in the loop • Pilot gate	• Framework first • Acceptable use and guardrails • Procurement and third-party risk • Continuous assurance	• Change by design • Psychological safety and trust • Skills uplift • Workforce and union considerations	• Control objectives • Platform strategy • Data handling rules • Intake and automation patterns • Endpoint-first controls

3.1 Governance

Governance is the anchor. As Plytas put it, “it all starts with governance.” Without a clear structure, AI either spreads in shadow form or stalls in indecision. Hillier noted, “you need board and senior staff buy-in, but the practical ideas will come from the people in the trenches.”

Good governance is more than process for its own sake. It is where data and stories surface, where accountability stays human. The EU AI Act points in the same direction, requiring impact assessments and documented oversight for high-risk uses.

3.2 Policy and compliance

If governance is the anchor, policy is the bridge between principle and practice. It tells staff what is allowed, what is not, and how risks are managed. Byrne put it plainly: “You’re better off having them in the tent than operating outside the tent.”

Examples show what this looks like in practice. Canada’s Directive on Automated Decision-Making requires agencies to complete Algorithmic Impact Assessments and publish public registries, making transparency part

of daily work. Hillier’s advice was simple: “Start with what you can do to ensure it is used safely, rather than saying you cannot use it.”

Genders urged balance in the matter: “Any attempt...to put on a relatively strong policy is going to be seen as draconian...the other side’s got to be about education.” Policy backed by education creates adoption. Policy alone drives workarounds.

3.3 People and culture

In the context of AI, people and culture mean more than training staff to use new tools. They cover the skills, confidence, leadership, and organizational climate that determine whether adoption succeeds or stalls. Technology does not implement itself. Instead, it is shaped by how people respond and whether the culture supports safe experimentation.

Norton reminded us that training “is not one-size-fits-all,” and that progress depends on champions and small wins that build confidence over time. But skills alone are not enough. Wade Hillier warned that “fear of public exposure gets in the way. That is not a good rationale for doing nothing.” Staff need psychological safety—the

assurance that AI use will be transparent, that human checks remain in place, and that accountability does not disappear into the system.

Culture, leadership, and clarity are what turn AI from a tentative experiment into a sustainable practice.

3.4 Technology and controls

Technology and controls form the guardrails that turn aspiration into reality. Regulators need to know the quality, provenance, and security of their data. Byrne's questions get to the heart: "What is the quality of that data? Is there a bias? Data sovereignty? What is the provenance of that data?"

Also, with vendors embedding AI by default, regulators need to understand and rationalize their platform choices. As Norton noted, "AI is just one service" within an ecosystem, and controls should fit into existing identity and device management.

Data handling rules, intake and automation patterns, and—critically—endpoint-first controls for shadow AI round out technology and controls considerations. The data underline why this matters. A Capgemini study found that only 21% of public-sector organizations have the data needed to effectively train AI models, while an EY survey reported that 45% cite inadequate data infrastructure as a barrier.^{14 15} Privacy and security also consistently emerge as the top concerns—flagged by 79% in the Capgemini research and 62% in the EY survey. For regulators, explicit guardrails on data and security are not optional; they are the single biggest determinant of whether AI use will build or erode trust.

From pillars to practice

Governance provides the anchor. Policy and compliance give shape. People and culture determine whether adoption takes hold. Technology and controls ensure it happens safely. Each pillar reinforces the others. Without governance, policy becomes shelfware. Without policy, culture drifts. Without culture, controls are ignored. And without controls, governance is hollow. Taken together, these pillars give regulators a way to move from ad-hoc experiments to sustainable, trusted use of AI.

¹⁴ Capgemini Research Institute, *Data Foundations for Government: From AI Ambition to Action*, 2025.

¹⁵ EY, *How data, analytics and AI in government can drive greater public value*, 2025.

4. Applications in regulatory practice

When it comes to regulatory practice, concrete examples of AI in action remain limited. This paper touched on certain narrow examples of targeted AI pilots, but dynamic, sweeping applications of AI are elusive. Asked if he had seen any compelling real-world deployments by regulators, Byrne was blunt: “Honestly, compelling regulator-led deployments are scarcely seen so far.” Other interviewees echoed the point: potential uses are often discussed, but proven cases are still rare. The absence of mature deployments, however, should not be mistaken for lack of relevance. Interviewees uniformly agreed that the potential for AI in regulatory contexts is vast—and that adoption is, in some form, inevitable.

This section therefore explores where AI might add value in regulatory practice, setting out both opportunities and cautions. Some uses—such as drafting, clustering, or validation—have appeared in pilot projects. Others remain possibilities that regulators are considering, but have yet to test at scale.

4.1 Licensing and registration

Licensing is often cited as a strong candidate for AI, given its reliance on standardized documents and repeatable workflows. While few regulators have formally deployed AI in this space, the logic of automation is compelling.

Standardize and validate

AI could be used to check routine documents for completeness and flag missing fields. Plourde suggested this is where AI can help “move files along, reducing impediments.”

Pre-approve the simple, review the complex

Routine renewals that meet set criteria could, in theory, be auto-drafted for human confirmation, leaving exceptions for full review.

Feedback loops from complaints

Analytics might eventually connect complaints data back into licensing processes, highlighting where risk-based verification should be strengthened.

Early evidence from financial supervision suggests AI can help triage filings and surface anomalies for human review, a pattern regulators can adapt for licensing workflows. The IMF notes that supervisors already

use AI for identity verification and anomaly detection across anti-money laundering and counter-financing of terrorism (AML/CFT) and reporting use cases—human-in-the-loop by design.^{16 17}

The opportunity is clear, but examples of regulators doing this today remain sparse.

4.2 Complaints and enquiries

Complaints and enquiries are high-volume, making them attractive for AI support. Here too, proven uses are limited, but the potential is easy to imagine.

Clustering and trends

AI could group complaints by theme and flag unusual surges. Hillier called complaints “the best data source. Bring them on and learn from them.”

Draft responses, human triage

Drafting replies to routine issues is possible, but sensitive matters must remain with people. Kelley noted that complaints can serve as a real-world test of risk-based approaches: “Complaints will tell us if our ‘trust the licensee’ approach works as we verify less and rely more on system checks.”

OECD’s latest Regulatory Policy Outlook highlights a shift toward data-driven, anticipatory regulatory delivery, pointing to AI-enabled text analysis and trend detection that can make complaint handling more proactive (while keeping sensitive triage with people).¹⁸

The likely near-term reality is AI surfacing patterns while staff continue to handle substantive responses.

4.3 Inspections and fieldwork

Inspections and fieldwork are often discussed as potential sites for AI adoption, but few field-ready examples exist.

Preparation

AI could create pre-visit packs summarizing history and risk cues, helping inspectors prepare more efficiently.

On-site tools

Mobile tools with transcription or evidence capture could lighten the administrative load, but they remain largely aspirational.

¹⁶ International Monetary Fund, *Artificial intelligence and the rise of regtech*, 29 October 2021.

¹⁷ International Monetary Fund, *Powering the digital economy: Opportunities and risks of artificial intelligence in finance*, Departmental Paper No. 21/24, October 2021.

¹⁸ Organization for Economic Co-operation and Development, *Regulating for the future: OECD regulatory policy outlook 2025*, 2025.

Post-visit summaries

Auto-drafted reports are possible, but always with human review and sign-off.

In supervisory contexts, SupTech pilots use machine-learning on regulatory returns to flag outliers much earlier—an approach regulators can mirror in inspections by pre-packaging risk cues before a site visit. Case examples from the UK show supervisors exploring ML to identify patterns in reporting and route scarce inspection capacity to higher-risk cases.¹⁹

The caution is that inspectors must document their rationale and avoid over-relying on prompts.

4.4 Investigations and enforcement

Investigations are where most interviewees drew the sharpest boundaries. The consensus was that AI might assist with background work, but determinations must remain human.

Assistive analysis

AI could be useful in searching documents, summarizing records, or highlighting links.

Human control

Decisions and enforcement actions must remain strictly human. Maintaining chain of custody and explainability is non-negotiable.

In high-risk domains like pharmaceuticals, industry pilots show AI assisting with submission drafting, pharmacovigilance signal detection, and quality documentation—with lifecycle validation and human oversight emphasized throughout. This reinforces your caution that AI should assist investigations, not determine outcomes.²⁰

Plourde cautioned: “Do not use AI to investigate; these are human, trauma-sensitive interactions... AI can transcribe only.”

4.5 Policy and guidance

Policy and guidance are text-heavy, making them a natural testbed for AI. Some regulators are experimenting here, though formal adoption remains limited.

Drafting with context

AI can draft guidance documents grounded in legislation, leaving experts to refine and approve.

McIntosh advised: “Utilize technology plus subject-matter experts—deploy AI co-pilots.”

Verification and approvals

Every draft must still be reviewed, versioned, and formally approved before publication.

Regulators can also borrow from Automated Regulatory Intelligence (ARI): using AI to monitor rule changes, classify obligations, and draft first-cut summaries for expert review—reducing manual “reg-watching” toil while improving currency and consistency of guidance.²¹

Here, the opportunity is clearer: regulators already use AI informally for drafting, but institutionalizing this safely requires strong verification processes.

4.6 High-risk domains: healthcare and justice

Beyond licensing and complaints, research shows certain sectors demand even greater assurance. Healthcare is one: while AI offers promise in diagnostics and administrative support, regulators must require validation standards and lifecycle monitoring of systems in use. Canada’s aforementioned directive and related federal guidance stress precisely this need for lifecycle oversight and public accountability, further detailed in the responsible use framework.

In law enforcement and justice, the risks are more fundamental: bias in facial recognition and predictive analytics has already drawn public concern. Documented uses in these domains underscore the need for stricter human oversight and transparency, with studies such as the European Commission’s *Ethics Guidelines for Trustworthy AI* highlighting the rights-based risks.

From possibilities to practice

The examples provided in this section are better seen as possibilities than established practice. Interviewees acknowledged that proven, regulator-led deployments remain rare, but the areas where AI could help are becoming clearer. Analysis of public-sector work shows that a significant share of time is concentrated in administrative tasks—meetings, scheduling, form processing, and record management—the very activities that underpin much of regulatory practice. These are natural candidates for early pilots, provided human review and accountability remain in place.

¹⁹ Regnology, *Future of SupTech: AI and machine learning in regulatory reporting*, 2023.

²⁰ IntuitionLabs, *AI and the future of regulatory affairs in the U.S. pharmaceutical industry*, May 2025.

²¹ CUBE, *Understanding automated regulatory intelligence*, September 2024.

By contrast, high-risk areas such as healthcare and justice demand deeper assurance, stronger standards, and visible human oversight. The role of regulators is to separate speculation from practice, prioritize early testing where value is high and risk is low, and build governance that can scale as concrete examples emerge.

Across licensing, complaints, inspections, investigations, and policy work, the most credible near-term gains come from AI assisting discovery, triage, and first-drafting—with people deciding. OECD’s 2025 Outlook points to anticipatory, data-driven regulatory delivery; the IMF documents concrete supervisory uses in AML/CFT and reporting; and ARI research shows how monitoring and summarizing obligations can be automated for expert review. Taken together, these strands validate an “assist, don’t decide” posture and suggest practical pilots that build capability without over-promising.



5. Practical pathways to get started

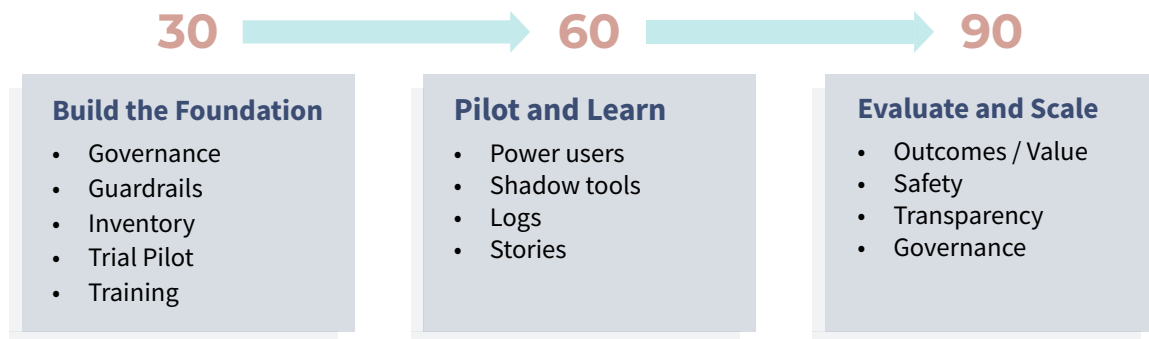
Every regulator interviewed for this paper agreed on one thing: the best way to move from talk to practice is to start small.

AI readiness will not come from writing policies alone or waiting for perfect examples. Instead, it will emerge from running controlled pilots, testing guardrails, and building confidence through measured steps. Norton described the pattern clearly: “start with a power-user cohort, provide training and support, run a three-month pilot, then review and if successful, expand.” McIntosh added that sometimes what is needed is simply “a focused four-hour workshop with the right people... come out with a plan.”

The 30–60–90 framework that follows is not a rigid recipe. It is a starting point—a way to give structure to the first three months of AI exploration, so that learning is captured and mistakes are contained.

5.1 The 30–60–90 plan

The first 90 days of AI adoption should be structured, visible, and deliberately narrow in scope. The goal is not to solve everything at once, but to create momentum, establish habits, and generate proof points.



Days 1–30: Build the foundation

- **Form the governance group.** Identify a small, cross-functional team that includes legal, IT/security, privacy, records, and frontline operations. Give it a clear mandate: make decisions, record them, and report back. At this stage, a light but explicit charter is enough—what matters is that someone owns the AI file, with clear sponsorship from executive leadership to signal priority and accountability.
- **Publish a one-page interim guardrails note.** Do not wait for a perfect policy. Staff need something to refer to right away. The one-pager should cover acceptable use, prohibited use, what to do with sensitive data, and when human approval is required. Pair it with a short explainer session so staff know it exists and how to apply it.
- **Stand up an AI-use inventory and begin discovery of shadow tools.** Ask every division what tools they are already using or experimenting with. Combine formal tools (e.g. Copilot, Salesforce) with shadow use (personal ChatGPT accounts, browser plugins). The first inventory won't be comprehensive, but it will make invisible use visible.

- **Choose a narrow internal pilot with clean data and motivated users.** Look for a process that is low-risk but high-friction—something that annoys staff but doesn't affect public rights. Meeting notes, case summaries, or internal drafting tasks are often a good start. Data availability is critical: if the data isn't clean or accessible, shelve that idea and pick another.
- **Deliver a one-hour crash course.** Cover basics of prompting, verification, privacy hygiene, and records obligations. The training doesn't need to be perfect—the act of showing staff that leadership is engaged matters as much as the content. Training refinements will come in future phases. One hour of training may lack comprehensiveness, but a simple course to bring staff up to speed
- **Pair training with policy visibility.** Many staff have had some exposure to AI basics, yet a significant share still isn't sure what the organization's rules are. A concise one-page guardrails note, paired with a short walk-through in month one, is one of the highest-leverage steps you can take to make safe use concrete.

Days 31–60: Pilot and learn

- **Run the pilot with a power-user cohort.** Start with a small group of trusted staff who are motivated to test and provide feedback. Within the group, establish a power-user cohort. Their feedback will surface both technical issues and cultural concerns.
- **Clamp shadow tools and provide sanctioned alternatives.** By this stage, you should know what unsanctioned tools are in use. Turn off the riskiest ones and offer safe alternatives. The message should be clear: you don't have to stop experimenting, but you do need to use approved channels.
- **Instrument logs to capture interactions and exceptions.** Do not rely on anecdotes. Build in basic monitoring so you can see usage patterns, where prompts fail, and what exceptions are triggered. This data will feed governance reviews and reassure executives that the pilot is controlled.
- **Collect value stories.** Ask participants to document time saved, pain points reduced, and lessons learned. These stories will be more powerful than metrics when you seek buy-in for scaling.
- **Refine the guardrails and prompts library.** Based on what you learn, update the one-page guardrails note and begin creating a library of sample prompts and verification checklists that staff can reuse.

Days 61–90: Evaluate and scale cautiously

- **Evaluate outcomes against entry and exit criteria.** Did the pilot achieve its goals? Were error rates acceptable? Were risks managed? Without meeting pre-defined criteria, the pilot should not expand. Treat a failed pilot as a lesson learned, not a setback.
- **Scale safely to the next team if results are positive.** If the pilot shows value and risks are controlled, extend it to a second group. Continue to document results and collect stories.
- **Publish transparency notes.** Communicate openly with staff and stakeholders about where AI is being used and where human oversight is applied. Transparency builds trust and reduces rumors.

Schedule quarterly governance reviews and continuous assurance checks. The pilot is not the end; it's the beginning of a habit. By day 90, governance reviews and monitoring should be on the calendar, turning experiments into ongoing practice.

5.2 Pilot design and measurement

Pilots succeed or fail on design. Define success up front in terms that boards and executives

recognize: throughput, timeliness, quality, rework, and satisfaction. Capture baselines before you begin. Include error-handling and rollback plans so experiments don't create crises.

Treat data quality as a precondition, not an afterthought. Most AI stumbles come from the plumbing—fragmented sources, unclear ownership, stale or inconsistent records—rather than the model itself. Make this a hard gate: “Data availability and quality: GREEN before launch.” If inputs aren't clean and accessible, pause or choose a different workflow; even the best-designed pilot will stall without sound data.

What “GREEN” looks like in practice:

- **Completeness & accuracy:** Required fields present, low error rates on a representative sample (e.g., $\geq 95\%$ field completeness; $\leq 2\%$ critical errors).
- **Consistency:** One “golden” source of truth for key entities; no conflicting values across systems.
- **Timeliness:** Data is current enough for the decision window (clearly defined freshness thresholds).
- **Provenance & permissions:** You know where the data came from, who owns it, and you have documented permission to use it for the pilot; sensitive fields are minimized or redacted.
- **Bias & representativeness:** The sample reflects real-world cases, including edge cases; you've checked for skew that could mislead outcomes.

For common applications and text-based productivity pilots (drafting minutes, document comparisons, summarization):

- **Scope:** Clear boundaries on what types of documents may and may not be used.
- **Redaction:** Sensitive or confidential information removed before use.
- **Review:** Human sign-off required for any drafted output before circulation.
- **Records:** Outputs logged and stored according to records/FOI obligations.
- **Transparency:** Clear disclosure where AI assistance was used.

Make a Pilot Data Pack a required deliverable before build:

- **Source inventory & ownership:** Systems, tables, APIs, data owners/stewards.
- **Schema & dictionary:** Field definitions, valid values, units, and business rules.
- **Profiling results:** Missingness, duplicates, outliers, freshness metrics, known quality issues.

- **Ground-truth set:** A small, curated set (e.g., 100–300 cases) with verified outcomes to evaluate draft/triage quality.
- **Redaction plan:** How PII/confidential fields are excluded or masked; retention & disposal rules
- **Acceptance thresholds:** Clear pass/fail criteria for data quality and coverage.

Two pragmatic tips:

- **Small clean beats big messy.** Start with a narrow, high-quality slice you can trust; widen later.
- **If retrieval is involved (e.g., internal document assistants),** verify document authority and recency, and log which sources are cited for every answer

Finally, wire data checks into the pilot: automated validation on ingest, drift monitors on key fields, and a short weekly review of exceptions. Good data makes measurement credible, keeps risk low, and turns a one-off pilot into a repeatable pattern.

5.3 Capability uplift program

Tools won't make you ready—capabilities will. In regulation, AI competence is less about mastering a specific product and more about building repeatable habits: safe prompting, rigorous verification, privacy hygiene, clear disclosure, and knowing when to escalate.

Capability uplift must run in parallel with pilots so people can use AI confidently, within boundaries, and with a shared language for quality.

A. Foundations for everyone (baseline in month one)

- Core concepts: what AI can/can't do; “assist, don't decide”; human-in-the-loop.
- Safe use: acceptable/prohibited uses, privacy hygiene, redaction, records/FOI obligations, disclosure language (“assisted by AI”).
- Prompting & verification: structure prompts; cite sources; check for gaps; how to escalate.
- Artifacts to ship: 1-page guardrails, “when in doubt escalate” card, redaction cheat sheet, verification checklist.

B. Role-based deepening (tracks by function)

- Case officers/licensing: intake triage, license renewals, criteria checks, standard letters; quality criteria (consistency, completeness).
- Investigations/enforcement: transcript prep, link analysis aids, chain-of-custody discipline, trauma-informed communication; strict decision boundaries.

- Policy/legal: reg-watch summaries, impact assessment drafting, version control, citation and provenance.
- Inspections/fieldwork: pre-visit packs, offline capture, rationale logging, bias traps; safety and retention rules.
- Communications: plain-language drafting, fact checks, citation standards, public disclosures.
- For each track: scenarios, red flags, approved prompt patterns, review checklist, escalation tree.

C. Enablement & support (make it easy to do the right thing)

- Champions network: 1–2 people per unit; “train-the-trainer” model; office hours and show-and-tell.
- Prompt & pattern library: vetted prompts, use-case templates, “before/after” examples; searchable and versioned.
- Help channel: a single place (Teams/Slack/SharePoint) for questions, tips, and recording FAQs.
- Safe sandboxes: non-production environments with synthetic/redacted data to practice safely.

D. Leadership & governance capabilities (sponsor the habits)

- Executive briefings: risk-tiering, decision rights, procurement clauses, privacy impact assessments, what to ask vendors.
- Manager playbook: how to approve a use, read an exception dashboard, run a lightweight post-incident review.
- Tabletop exercises: shadow-AI leak, mis-messaging, embedded-feature “on by default”; who speaks in the first hour.

E. Rhythm, measurement, and reinforcement

- 30–60–90 cadence:
 - 30 days: foundations + guardrails walk-through; champions named; prompt library v1.
 - 60 days: role tracks live; office hours; first “value stories” published.
 - 90 days: refresh guardrails; incorporate lessons from pilot; add modules based on exceptions seen.
- Metrics that matter: participation/completion, quiz pass rates, policy-exception rate, QA scores on AI-assisted drafts, time-to-first-value, help-channel response time.
- Recognition: highlight value stories; “prompt of the month”; badges for champions.

F. Culture, safety, and inclusion

- Psychological safety: normalize “show your work,” share prompts, and log corrections; blameless post-mortems.
- Accessibility: plain-language materials; multiple formats (video, 10-minute micro-modules, step-by-step guides).
- Workforce engagement: brief unions/staff associations early; emphasize augmentation and redeployment over replacement.

Deliverables checklist (lightweight, practical)

- Guardrails one-pager (v1), verification checklist (v1), redaction cheat sheet, disclosure language bank.
- Role cards (decision boundaries per role), escalation tree, prompt/pattern library (searchable, versioned).
- Sandbox access guide with do/don't examples; short “manager's approval” form for new uses.

Bottom line: capability turns pilots into practice. Group training by **who does what**, back it with simple artifacts, and build a steady rhythm of learning, measurement, and recognition. That's how you move from awareness to embedded habit—safely, visibly, and at pace.

5.4 Procurement fast lanes

Procurement often slows down innovation, but with AI now **embedded in mainstream SaaS and platforms**, regulators can create *fast lanes* that enable small, safe trials without waiving safeguards. The aim is to move quickly on **low-risk, reversible pilots** while baking in the controls you'll need at scale. Ontario's *Responsible Use of AI* directive shows how to embed expectations in procurement; U.S. OMB guidance likewise pushes agencies to integrate AI risk management into contracts. Use those as anchors, then make the following your house playbook.

A. When to use a fast lane (eligibility)

Choose pilots that are:

- **Low risk:** internal-only use, “assist/draft—not decide,” no rights-affecting outcomes.
- **Reversible:** clear kill-switch, no lock-in, data export supported.
- **Data-mature:** clean, non-sensitive inputs or redacted/synthetic datasets.
- **Scoped:** ≤ 90 days, capped users, clear success/exit criteria, sandbox or test tenant.

B. Contract essentials (must-have clauses)

1. Data handling & residency

- Data stays in specified regions; no vendor training on your prompts or outputs (“no train on customer data”).
- Segregation of tenant data; clear data ownership and IP.
- Deletion/export SLAs at pilot end.

2. Model transparency & change control

- Versioning and release notes for models; 30-day advance notice for material changes.
- Opt-out of auto-enabled features; ability to roll back to prior version if changes raise risk.
- Disclosure of model type (e.g., provider, fine-tuned vs base), and inference location.

3. Built-in guardrails

- Configuration for draft-only outputs; human-in-the-loop checkpoints.
- Prompt/response logging under your control; configurable retention.
- DLP compatibility; upload/connector controls; content filters and jailbreak mitigations.

4. Security & assurance

- SSO/SAML required; least-privilege roles; admin telemetry.
- Audit rights; independent assurance (e.g., SOC 2 Type II/ISO 27001).
- Incident reporting within a defined window (e.g., 72 hours); sub-processor register.
- AI governance alignment: vendor states how they map to ISO/IEC 42001 and NIST AI RMF; provide model cards or equivalent documentation.

5. Kill-switch & portability

- Tenant-level disable control you can invoke unilaterally.
- Data return format specified; migration support if terminated.

6. Records & transparency

- Exportable logs for prompts, responses, user, timestamp, model version.
- Ability to surface disclosures in public-facing outputs (“AI-assisted”) when required.

7. Evaluation & testing rights

- Right to benchmark with synthetic/test data; cooperation on bias and quality tests relevant to the pilot.

C. Pilot commercial terms (keep it light, safe, and short)

- Short-form SOW + DPA; cap spend and seats; no auto-renew.
- Termination for convenience with minimal notice.
- Vendor provides named contact, change log, and quarterly review.

D. Quick vendor due-diligence screen (10 questions)

1. What model(s) power the feature? Who runs inference, and where?
2. Do you train or fine-tune on our data/prompts? Default stance?
3. How do you prevent prompt leakage across tenants?
4. Can we disable embedded AI features at tenant/feature level?
5. What default safeguards (toxicity, PII, jailbreak) are enforced? Configurable?
6. How are model updates communicated? Is there an opt-out/rollback path?
7. What logs can we export (prompts/responses/versions)? Retention settings?
8. Do you support SSO, role-based access, and admin telemetry?
9. What third-party assurance do you hold (SOC 2/ISO 27001; AI governance statement vs ISO 42001/NIST AI RMF)?
10. List sub-processors and their locations; provide incident response playbook.

E. Operating model with vendors (don't "procure by accident")

- Require a feature-toggle register: what AI features exist, and who approved them.
- Route all material model changes through a lightweight risk check (privacy, security, records, comms).
- Hold a quarterly vendor review (usage, incidents, upcoming changes, assurance status).
- Tie procurement to governance artefacts: pilot decision log, AI-use inventory, risk register.

F. One-page checklist (to staple to every AI SOW)

- Scope: assistive only, no automated determinations; users capped; time-boxed.
- Data: residency stated; PII plan; "Data quality = GREEN" verified.
- Controls: SSO, logging, draft-only, kill-switch, opt-out of auto-features.
- Legal: DPA, audit rights, update notices, rollback, deletion/export SLAs.
- Assurance: mapping to ISO 42001/NIST AI RMF; incident window; sub-processor list.
- Comms: disclosure language agreed; transparency note template ready.

Why this matters now

Embedded AI arrives by default; without these terms and an operating rhythm, you inherit risk you didn't plan for. A fast-lane approach lets you move quickly on low-risk pilots and prove to boards and the public that controls, transparency, and reversibility are in place.

On practical pathways

Practical progress comes from pilots that are narrow, measurable, and well-governed. The first 90 days are about building momentum: standing up governance, publishing guardrails, running one pilot, and capturing lessons. Start with areas where data is mature, staff are motivated, and risks are low. Pair training with visible policy so staff know both how to use AI and what the boundaries are. And put in place procurement and governance rhythms that can scale as pilots turn into practice.

6. Consequences of inaction

Regulators face a choice: start experimenting now with guardrails in place, or risk being overtaken by events. Every interviewee emphasized that doing nothing is not neutral. AI is already embedded in platforms, staff are already experimenting, and licensees are already adopting it. Failure to respond leaves regulators exposed on multiple fronts—from security breaches to reputational damage.

Byrne's calculus: inaction carries greater risk than careful, bounded experimentation. Norton was blunter still, suggesting that regulators who deploy AI without controls are setting themselves up for problems. This section lays out what inaction looks like: the risks, the scenarios regulators may encounter, and the importance of trust and communications in managing the fallout.

6.1 Risk taxonomy

The risks of inaction are not hypothetical. They are visible today in how AI tools are being embedded, used, and misused.

Privacy and confidentiality exposure

Employees pasting sensitive complaints or records into unsanctioned tools can cause instant breaches.

Inaccurate outputs

Drafted content that slips through unchecked can harm individuals or create inconsistent regulatory outcomes.

Bias and fairness concerns

Algorithms may reinforce inequalities, especially in justice, healthcare, or licensing.

Loss of public trust

Opaque or inconsistent use of AI risks eroding confidence in regulators.

Regulatory lag

Falling behind fast-moving markets makes enforcement reactive rather than proactive

Staff workarounds

When guardrails don't exist, staff find their own paths, bypassing controls and increasing risk.

Surveys reinforce these concerns. Privacy and security are consistently cited as the top barriers to AI adoption—79% in a Capgemini study, and 62% in an EY survey. For regulators, those barriers are not abstract; they are the very risks that come with delay.

6.2 Scenario snapshots

Scenario-based thinking helps regulators plan for incidents before they happen. Three common risks illustrate the point:

- **Shadow tool leak.** An employee pastes a sensitive complaint into an unsanctioned chatbot. *Prevention:* training and endpoint controls. *Detection:* egress monitoring. *Response:* disclosure to affected parties and remediation.
- **Embedded AI turns on.** A SaaS platform quietly enables an AI feature that ingests personal data. *Prevention:* contract clauses and vendor change notifications. *Detection:* vendor notices, monitoring logs. *Response:* disabling the feature and reviewing data flows.
- **Automated mis-messaging.** A draft response generated by AI is sent without human review. *Prevention:* configure tools to output drafts only; set clear approval gates. *Detection:* spot checks. *Response:* correction and transparent communication.

These scenarios underline why pilot gates, procurement clauses, and training are so important. Incidents don't happen because regulators "chose AI"—they happen when regulators fail to prepare for what AI is already doing in their environment.

6.3 Public trust and communications

Public trust is the regulator's true currency. When something goes wrong—and something will—the question is whether you can explain it quickly, clearly, and honestly.

Transparency

Be explicit about where AI is used and where human oversight applies. Publish short, plain-language notes for the public.

Preparedness

Have a "first hour" communications plan for incidents. Who speaks, what they say, and how fast they say it matters as much as the technical fix.

Education

Ensure staff and stakeholders understand both the potential and the risks. Hines warned that "the biggest risks are releasing private information; take an inventory before it's too late to control it." Kelley added that these

risks exist even without AI: “Clarity and education are your best defenses.”

Inaction is, in effect, an abdication of responsibility. Regulators don’t get to choose whether AI shows up in their environment—it already has. What they do control is whether its risks are managed transparently and its benefits channeled toward public value.

The cost of doing nothing

If regulators ignore AI, the risks don’t disappear—they grow.

Privacy breaches

Sensitive information slips into unsanctioned tools, creating instant compliance failures.

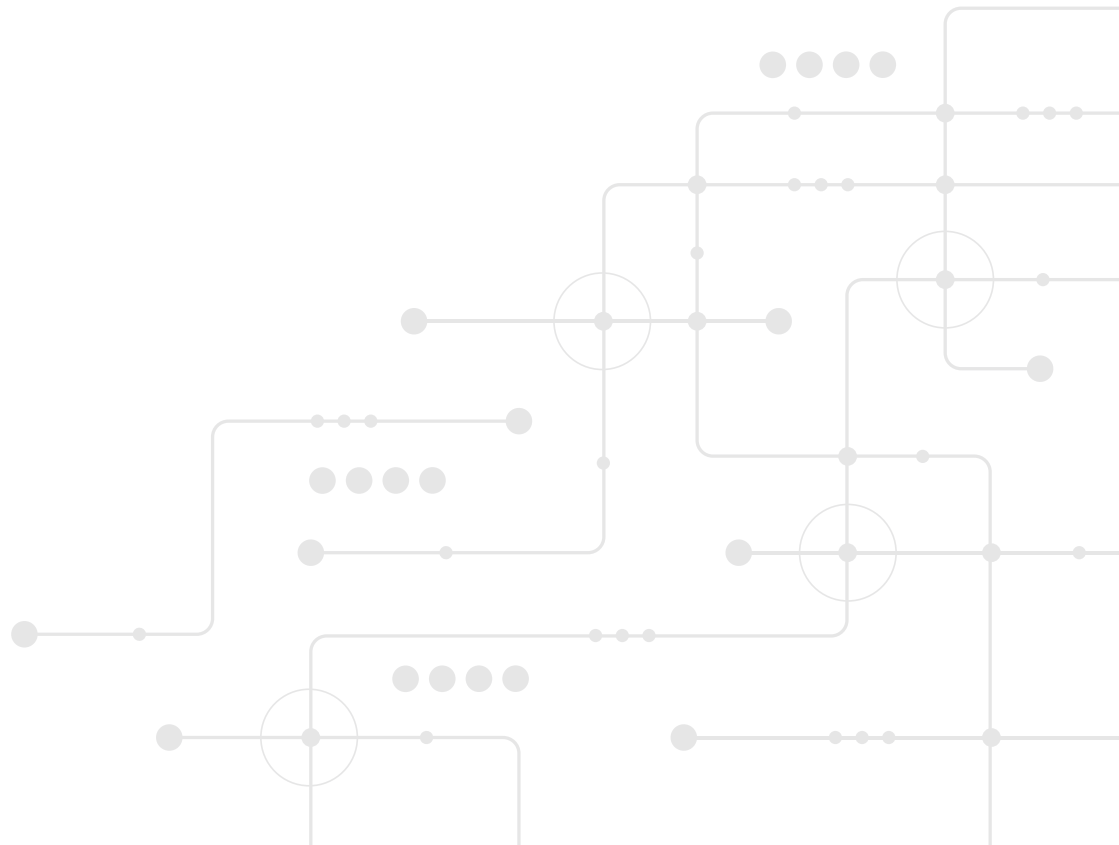
Loss of trust

Opaque or inconsistent AI use erodes public confidence in the regulator.

Falling behind

Markets adopt AI faster than regulators, leaving oversight reactive instead of proactive.

Doing nothing is not neutral—it is itself a high-risk choice.



7. Recommendations and next steps

By this point, the pillars are clear and the risks well understood. What matters is translating them into action.

Regulators told us that moving forward requires visible leadership, controlled experimentation, and credible signals of readiness. McIntosh urged executive sponsorship: “Go straight to the board and CEO...you need that level of buy-in to make this successful.” Plytas emphasized credibility: “Certification and audits under ISO and NIST build trust and signal readiness.”

This section sets out a practical roadmap: what to do in the next 30 days, what to expand in the first 100 days, where to aim after a year, and how to collaborate across borders.

7.1 Immediate actions—next 30 days

The first month is about getting organized and signaling intent.

- **Form the governance group and publish interim guardrails.** Even a light-touch committee and a one-page acceptable use note send a clear message: AI is being taken seriously, and staff have guidance.
- **Start the AI-use inventory and map shadow AI.** Visibility matters. You cannot govern what you cannot see.
- **Choose one pilot with clean data and clear value.** Pick a low-risk workflow where data is available and quality is high. Define entry/exit criteria so you know what success looks like.
- **Schedule a crash course and appoint champions.** A one-hour staff session covering prompting, verification, and privacy hygiene is enough to start. Identify early champions who can answer questions and reinforce safe use.

7.2 Next 100 days

With foundations in place, the focus shifts to proving value and tightening controls.

- **Run and evaluate the pilot.** If it works, scale cautiously to a second team; if not, document lessons learned.
- **Adopt basic continuous assurance.** Move beyond paper audits. Instrument logs, set up exception dashboards, and hold periodic reviews so issues are caught in real time.

- **Implement a procurement fast lane.** Use standard AI clauses to reduce delays and manage risks: model-update notices, opt-outs, kill switches, audit rights.
- **Publish a transparency note.** Tell staff and stakeholders where AI is being used and where humans review. Transparency builds trust and reduces rumors.

7.3 One-year roadmap

At the one-year mark, AI readiness should be embedded into the organization’s DNA.

- **Integrate AI oversight into enterprise risk management.** Make AI part of internal audit and risk registers.
- **Expand training and sustain learning.** Grow role-based curricula, maintain a shared prompts library, and formalize a community of practice.
- **Extend use cases cautiously.** Move beyond internal productivity to inspections and customer contact, but keep humans firmly in the loop.
- **Formalize collaboration with peers.** Share templates, model clauses, and lessons across agencies.

This is also the stage to invest in capability building. Multi-study consensus highlights the need for an uplift program that blends AI literacy, leadership development, and change management. Building this capacity is as important as the technology itself.

7.4 Collaborate across jurisdictions

No regulator can go it alone. International coordination avoids duplication, lifts standards, and reassures the public.

- **Risk-based by design.** Borrow from the EU’s approach: tier uses by risk (unacceptable, high, limited, minimal) and map your internal controls accordingly. A one-page mapping tool can help executives right-size governance.
- **Borrow boldly.** Canada’s mandated impact assessments and registries, and Australia’s lifecycle standards and GovAI collaboration platform, offer model clauses and templates regulators can adapt.

- **Explore new governance models.** Some experts propose “regulatory markets”—licensing private assurance providers to monitor AI developers under public oversight. This is not mainstream yet, but worth watching.
- **Join international efforts.** Polling shows strong public preference (41%) for international AI rules, especially in sensitive domains. Aligning with peers builds legitimacy.

Building trust through certification

- **ISO 42001** and **NIST AI RMF** offer credible benchmarks regulators can align with.
- Certification and audits against these standards send strong signals to boards, governments, and the public.
- *“Certification and audits under ISO and NIST build trust and signal readiness.”*
—George Plytas

Recommendations: closing note

Readiness is not a one-off project. It is a journey of early pilots, measured scaling, and continuous assurance. By starting with small, controlled steps and borrowing from proven frameworks internationally, regulators can act quickly without losing public trust. The goal is not perfection but progress—moving from ambition to action in ways that can be explained, defended, and sustained.



Conclusion—and looking ahead

This paper has mapped a readiness blueprint for regulators in the age of AI. We began by setting out why AI readiness matters, showing that while adoption is still uneven, the pace of change means inaction carries its own risks. We then built a shared understanding of the AI landscape: the tools regulators are most likely to encounter, where they might apply, and why principle-based frameworks such as ISO 42001, NIST’s AI RMF, and the EU AI Act provide practical anchors.

From there, we explored the four core pillars of readiness—governance, policy and compliance, people and culture, and technology and controls—illustrating how each reinforces the others. We looked at potential applications across regulatory practice, from licensing and complaints to inspections and investigations, and highlighted where pilots might start and where caution must be strongest. We outlined practical pathways, including a 30–60–90 plan, capability uplift, procurement fast lanes, and a one-year roadmap for embedding AI oversight. We also drew attention to the consequences of inaction: privacy breaches, trust erosion, regulatory lag, and the inevitability of shadow AI. And we closed with recommendations: start small, scale cautiously, learn in public, and collaborate across jurisdictions.

This blueprint is not a technology plan; it is a capability plan. AI will keep changing. As Genders framed it, capability is diffusing fast—good and bad alike—which is why controls and transparency matter. The practical response is measured: clear roles and gates, human-in-the-loop decisions, contracts that anticipate model changes, and continuous assurance that surfaces exceptions early. It is also cultural: education over prohibition, learning by doing, and sharing prompts and patterns that work. And above all, it is about trust. Trust is a regulator’s currency, as Byrne noted. Spend it wisely—by being open about where you use AI, careful about what it can and cannot do, and relentless about learning in public.

Contributors

Paul Byrne

Executive Director of Regulatory Operations & Support Services, Irish Medical Council; CLEAR President-Elect

Paul Byrne serves as Executive Director of Regulatory Operations & Support Services at the Irish Medical Council, where he oversees key regulatory functions including registration, education, professional competence, and performance reporting. He brings over 19 years of public policy experience with a regulatory focus, having previously held senior roles at CORU (Ireland's social care regulator) and in energy regulation. At CLEAR, Paul serves as President-Elect, chairs the Technology and Innovation Task Force, and is Vice Chair of the Regulatory Agency Administration Committee.

Wade Hillier

Regulatory Affairs Consultant; Former Deputy Registrar

Wade Hillier is currently a Regulatory Affairs Consultant who previously served as Deputy Registrar at the Retirement Homes Regulatory Authority (RHRA) in Ontario. His extensive experience includes significant time with both the Retirement Homes Regulatory where he was responsible for core regulatory functions and the College of Physicians and Surgeons of Ontario, where he worked in quality management and government programs including methadone oversight.

George Plytas

Senior IT/AI Governance Expert; CEO of Cyntry

George Plytas is a senior IT and AI governance expert and CEO of Cyntry, where he provides managed cybersecurity services, IT audits, and compliance attestations for small to medium-sized companies. He has held senior management roles across industries including insurance, payment processing, and automotive, developing security programs that elevated organizations to market-leading positions. Specialising in PCI-DSS, SOC1, SOC2, HIPAA, and ISO27001 compliance, he brings extensive knowledge in securing data centers and enterprise systems. Plytas's previous roles include Acting CISO at Moneris Solutions, Head of Security at CAA SCO, and senior consulting positions in information security and risk management.

Bernard Plourde

Registrar, College of Allied Health Professionals of PEI (CAHPPEI)

Bernard Plourde serves as Registrar of the College of Allied Health Professionals of Prince Edward Island (CAHPPEI), which regulates Medical Laboratory Technology, Medical Radiation Technology, and Respiratory Therapy in PEI. Notably, Plourde serves as registrar for multiple professional colleges in PEI, including the College of Paramedics, College of Counselling Therapy, and College of Occupational Therapists. He also serves on the Canadian Alliance of Medical Laboratory Professionals Regulators (CAMLPR) Board of Directors. His approach emphasizes transparency, accountability, and working collaboratively to ensure public safety across multiple healthcare professions.

Heather A. Kelley

Director of Operations, New Hampshire Office of Professional Licensure & Certification (OPLC)

Heather A. Kelley serves as Director of Operations for the New Hampshire Office of Professional Licensure & Certification (OPLC), bringing two decades of experience in local, state, and federal government service. In her role, she oversees operational functions for New Hampshire's professional licensing and certification programs. Kelley has been recognized for her expertise in organizational development, risk management, finance, and compliance monitoring, making her a key leader in modernizing professional regulatory services in the Granite State.

Andrew Norton

Chief Information Officer & Director, Business Operations, Law Society of Alberta

Andrew Norton serves as Chief Information Officer & Director of Business Operations at the Law Society of Alberta, a role he was appointed to in 2019. Originally from England, Norton joined the Law Society in 2011 as Director of Business Technology and has steadily expanded his influence across organizational operations. He is responsible for leading the organization through business operational changes, overseeing technology implementation, and managing departments including Membership and Information Management. Norton brings an MBA and extensive experience in Information Technology, Project Management, and Business Improvement from his previous work with UK emergency services and local government.

Ronne Hines

Former Director, Division of Professions & Occupations, Colorado DORA; Former CLEAR President

Ronne Hines is a regulatory consultant who previously served as Director of the Division of Professions and Occupations (DPO) at the Colorado Department of Regulatory Agencies (DORA), overseeing licensing for more than 50 professions and 500,000 licensees. She served as CLEAR's Past President and was recognized as a visionary leader in healthcare regulation and consumer protection. With 15 years in regulatory administration and a legal background, Hines transitioned from private practice to become an accomplished senior leader, focusing on governance, risk management, strategic planning, and organizational change management. She holds a Juris Doctor from the University of Denver and a Master's in Public Administration from the University of Colorado Denver.

Steve Genders

Information Security Architect, tekAssembly; Former VP Information Security

Steve Genders currently serves as Information Security Architect at tekAssembly, having previously held the position of VP of Information Systems and Security at a regulatory technology company for over four years. He brings more than two decades of experience in information security from major technology and financial firms, including IBM and BMO Financial Group. Genders holds multiple cybersecurity certifications including CISSP, CCSP, and CISM, and has been instrumental in building security programs from the ground up, implementing cloud security strategies, and advancing technological innovation in cybersecurity.

Joe McIntosh

Senior Consultant; Former CIO, State of Oklahoma

Joe McIntosh is a Senior Consultant specializing in innovative technology solutions for state, local, and education entities. As the former CIO for the State of Oklahoma, he spearheaded the state's digital transformation, enhancing efficiency, citizen services, and data-driven decision-making. McIntosh brings over two decades of experience spanning public and private sectors, with expertise in architecting modern IT infrastructures, developing statewide data strategies, and implementing cutting-edge cloud and AI solutions. He holds an MBA in e-Commerce and multiple certifications in project management and agile practices.

Authors



Andre Forget
CEO, GRCS Partners Inc.



Paul Leavoy
Regulatory Writer & Researcher

As regulators face rising cyber threats, evolving legislative requirements, and growing expectations to safeguard public trust, building digital resilience is essential.

GRCS Partners works with regulatory bodies worldwide to provide guidance on cybersecurity, privacy, compliance, and AI governance, helping agencies navigate complex challenges with confidence.

For additional resources and information, visit grcspartners.com.

About GRCS Partners

GRCS Partners is an advisory firm with deep expertise in the regulatory sector, specializing in governance, risk, compliance, and security. We support regulatory authorities and regulated organizations with tailored programs, including cybersecurity assessments, AI governance, third-party risk management, and alignment with leading standards and frameworks such as SOC 2, ISO/IEC 27001, ISO/IEC 42001, HIPAA, PCI DSS, GDPR, the NIST Cybersecurity Framework (CSF), and the NIST Risk Management Framework (RMF). Through evidence-based insights and executive advisory services, we help leaders strengthen oversight, safeguard sensitive information, and maintain public trust.