

Beyond Tech

Information Technology, Security, and Privacy

ISO STANDARDS AND BEYOND

Understanding the Roles
of IT, Security, and Privacy

Quantum Computing and
Cryptographic Risk: Are
Organizations Ready?

The Competitive Edge:
Integrating Information
Security and AI Governance

Turning ISO/IEC 27001,
NIS2, and DORA into Real
Organizational Resilience

06



The Expert
Your Phishing Awareness Training Is Making People Worse at Spotting Phishing

10
The Expert
AI Governance Starts before Deployment: A Cybersecurity Perspective

20



Leadership
The Competitive Edge: Integrating Information Security and AI Governance

26
Innovation
After Mythos: What's Actually Shifting in Cybersecurity?

36



Business & Leisure
The Eternal City – A Glimpse of Rome

46
Opinion
Turning ISO/IEC 27001, NIS2, and DORA into Real Organizational Resilience

50
The Expert
Why Every Executive Should Understand Technology Risk?

56

Innovation

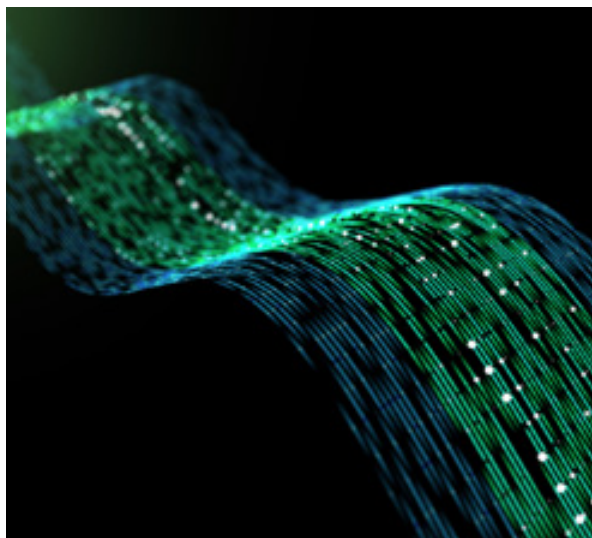
Privacy as a Foundation for
Innovation: Protecting Health
Data in the Digital Age

62

The Standard

How Quantum
Computers Work

72



Technology

Three Disciplines, One Mission:
Understanding the Roles of IT,
Security, and Privacy

78

Q&A from Webinar

The AI Governance Illusion:
Why ISO/IEC 42001 and CAIP
Matter More Than Ever

96

The Expert

Quantum Computing and
Cryptographic Risk: Are
Organizations Ready?

102



Work-Life-Balance

The Executive Resilience
Imperative: Work-Life Balance
in an Era of Constant Risk

110

Technology

Code is the New Compliance:
From Checkbox to Continuous

118



The Expert

Silence by Design: Why Good
People Stay Quiet About Breaches

128

The Expert

Risk-Informed Decision
Making: Helping Leaders
Navigate Digital Uncertainty



**“The greatest achievement
how it changes life, bu**

-Satya Nadella,

“

ment of technology is not
ut how it improves it”

CEO of Microsoft

THE EXPERT

Your Phishing Awareness Training Is Making People Worse at Spotting Phishing

Picture the surveillance audit. The awareness program is presented with justifiable pride: 100% completion of the annual security training, quarterly phishing simulations, a click rate falling steadily for three years. The auditor notes conformity with Clause 7.3 and moves on.

Everyone in the room believes the organization has become measurably harder to phish.

In 2025, that belief was finally tested with the rigor ISO management systems are supposed to embody, and it did not survive.

What the Evidence Actually Shows

The pivotal study was a randomized controlled trial published at IEEE Security and Privacy 2025, covering 19,789 employees of UC San Diego Health over eight months. Its design matters for anyone who thinks in terms of audit evidence: random assignment, control groups, and real behavioral outcomes rather than self-reported confidence.

The findings, in ascending order of discomfort. Annual awareness training showed no significant relationship between training recency and simulation failure; employees trained weeks earlier failed at the same rate as those a year overdue. Embedded post-click training produced an absolute improvement of under two percentage points, explained largely by engagement data showing a third to a half of employees closed the training page in zero seconds. And employees who completed multiple rounds of static embedded training exhibited 18.5% higher failure odds per additional session.

The intervention deployed to reduce the risk was, for a substantial population, associated with increasing it.

This corroborates rather than contradicts the earlier literature. An ETH Zurich field study of 14,733 employees over 15 months found the trained group subsequently clicked more than the untrained control, and identified a plausible mechanism: 43% of employees reported feeling safer after completing the training. The training delivered a sense of protection instead of protection, which is the definition of a control that has become a liability. A 2025 replication at a US financial services firm, with lure difficulty controlled using the NIST Phish Scale, found no significant training effect on either clicking or reporting.

For compliance professionals, the correct reading is precise: the evidence does not say awareness is worthless. It says the dominant delivery model, completion-tracked modules plus simulation click rates, produces activity without measurable effectiveness, and sometimes negative effectiveness.

This Is a Clause 9.1 Problem Wearing a Clause 7.3 Badge

ISO/IEC 27001 never asked for phishing simulations. Clause 7.3 requires that persons be aware of the information security policy, their contribution to the effectiveness of the ISMS, and the implications of nonconformity.



An internal auditor who accepts completion percentages and click rates as effectiveness evidence is, in ISO terms, accepting nonconforming measurement. The nonconformity just happens to be popular.

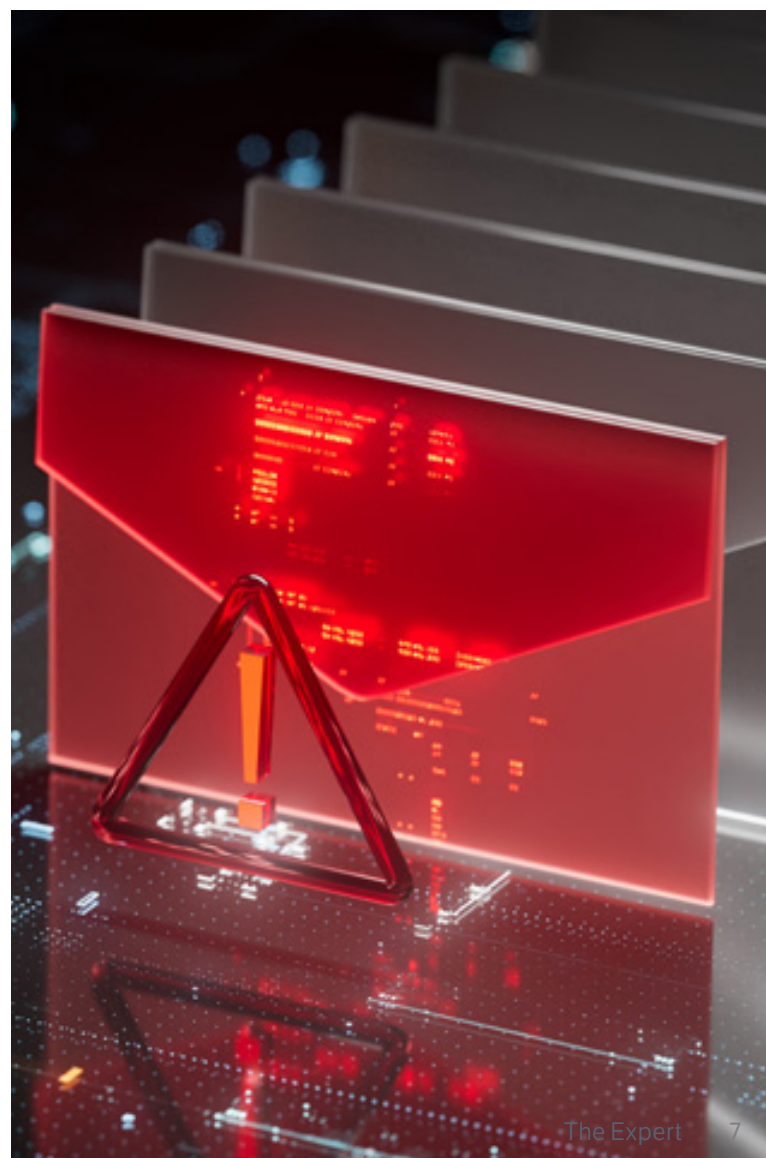
What Effective Looks Like, in the Language of the Standard

The same research literature that dismantles the standard model points clearly at what to build instead, and each element maps onto the ISMS rather than around it.

Measure reporting, not clicking. The ETH study found employee reporting operationally viable as a detection channel, with 68% accuracy and a manual triage load of roughly 1.5 emails per day, and the 2025 replication found that in over a third of campaigns someone reported the phish before anyone clicked. Reporting rate and time-to-first-report are genuine effectiveness metrics for Clause 9.1, and they connect awareness directly to the incident management controls (A.5.24 through A.5.26) where its value is realized.

Annex A control 6.3 requires awareness, education and training that is relevant to the person's role, with the supporting guidance in ISO/IEC 27002 emphasizing that awareness activity should be evaluated for effectiveness. And Clause 9.1 requires the organization to determine what needs monitoring, and to evaluate whether the ISMS achieves its intended results.

Now hold the standard program against those requirements. A training completion rate is a record that an activity occurred, which satisfies the documentation habit but says nothing about awareness outcomes; the zero-second engagement data shows exactly how wide that gap can be. A falling click rate looks like an effectiveness measure but fails basic measurement validity: in the UCSD trial, failure rates across lure templates ranged from roughly 2% to over 30%, meaning the metric responds more strongly to template selection than to workforce capability. A program can show three years of improving click rates while its people become no harder to deceive, and vendor before-and-after curves, constructed without control groups, should be treated with the same skepticism an auditor would apply to any supplier's uncontrolled performance claim.



A workforce that reports fast is a functioning detective control; a workforce with a low click rate on easy templates is a statistic.

Prefer controls that do not require vigilance. The UCSD authors' recommendation was not better training but phishing-resistant authentication and password managers. That is Annex A 8.5 territory, and the current threat environment reinforces it: adversary-in-the-middle kits now bypass one-time-code MFA at industrial scale, while FIDO-based passkeys remove the phishable credential entirely, with large-scale deployment data showing usability gains alongside the security ones. Risk treatment that shifts load from human judgment to technical control is simply proportionate treatment based on the evidence available.

Keep awareness, change its form. The follow-up ETH research published at ACM CCS 2024 found that brief periodic reminders performed as well as full training content, and that repeat clickers were failing through inattention rather than ignorance, which no module fixes. Short, frequent, blameless nudges satisfy the intent of Clause 7.3 better than the annual module, at a fraction of the cost, and without teaching employees that security is a trick played on them by their own organization. The leadership dimension of awareness culture matters more here than any content library: people report quickly when reporting is safe.

Retain simulations only as measurement. Used sparingly, unannounced to no one's punishment, and analyzed with difficulty held constant, simulations can measure whether reporting works and whether technical controls catch what humans miss. Used as training theatre, the evidence says they produce false confidence and audit-friendly numbers.

The question for your next management review

The uncomfortable summary is that a widely deployed control, purchased in good faith and documented impeccably, has failed its effectiveness evaluation in the largest trials ever run, and most ISMS documentation has not yet noticed. There is a certain irony available to organizations that pride themselves on investing in security awareness: the investment is sound; the instrument is not. The fix begins with one question in the next management review, asked exactly as Clause 9.1 intends: what evidence do we hold that our awareness program changes behavior, and would that evidence survive the scrutiny we apply to every other control? If the honest answer is a completion rate and a click-rate chart, the program has been auditing its own activity, and the organization deserves better data than that.

Awareness is still a requirement. Effectiveness always was.

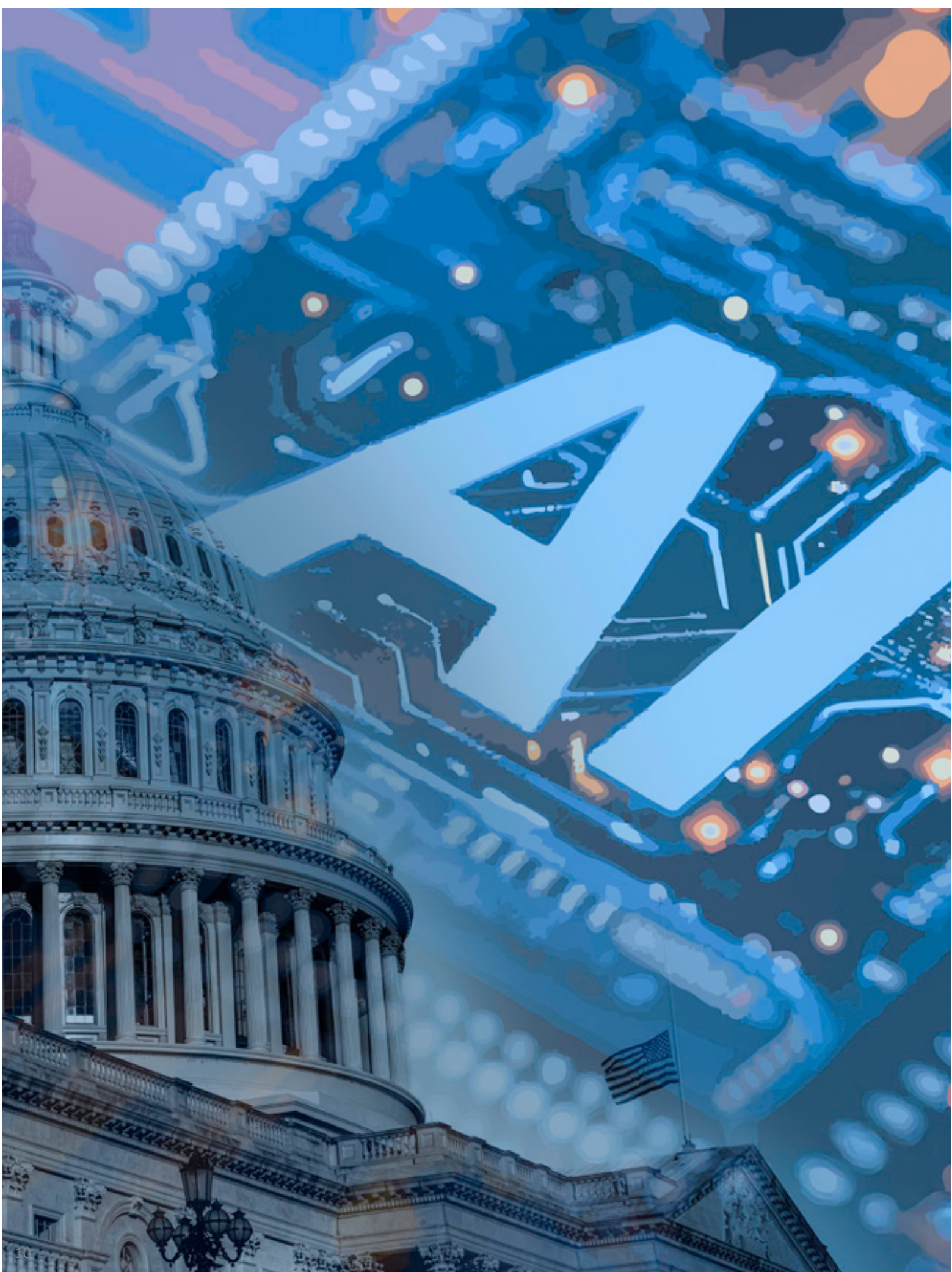




Nahla Davies

Software Developer

Nahla Davies is a software developer and tech writer. Before devoting her work full-time to technical writing, she managed — among other intriguing things — to serve as a lead programmer at an Inc. 5,000 experiential branding organization whose clients include Samsung, Time Warner, Netflix, and Sony. You can reach her at: nahlawrites.com



AI Governance Starts before Deployment: A Cybersecurity Perspective

Many organizations discover the need for artificial intelligence (AI) governance only when an AI solution is nearly ready to go live.

By then, the most important decisions have often already been made. A vendor has been selected, data has been shared, proof-of-concept work has begun, employees have tested external tools, and architecture choices have quietly shaped the future risk profile of the system.

That is the central problem. AI governance is often treated as a deployment activity, when in reality, it should begin with the first business decision to explore AI.

This matters because AI does not behave like an ordinary software project. It depends on data quality, third-party models, cloud services, prompts, APIs, training practices, and human decisions that may be difficult to reverse later. A model can be monitored after deployment, but confidential data uploaded during an early pilot cannot always be retrieved. A weak supplier agreement can be renegotiated, but only after the organization has inherited unnecessary risk. A poorly governed experiment can become a business-critical tool before anyone formally approves it.

From a cybersecurity perspective, the lesson is clear: the earlier governance is introduced, the more effective it becomes. Security teams add the greatest value not when they act as the final approval gate, but when they help shape business, procurement, data, and design decisions before risks become embedded.

Governance Begins with the Decision to Adopt AI

Every AI initiative should start with a business question, not a model. What problem is the organization trying to solve? Why is AI the right approach? What information will the system process? Who will rely on its outputs? What would happen if those outputs were wrong, manipulated, or unavailable?

These questions may appear simple, but they define the foundation of governance. They determine whether the initiative is low risk, business critical, privacy sensitive, supplier dependent, or potentially subject to regulatory obligations.

Too often, organizations ask these questions after a solution has already been chosen. At that stage, cybersecurity teams are asked to review architecture, evaluate controls, or approve deployment. Those activities remain important, but they cannot fully correct earlier decisions made without sufficient oversight.

Consider a procurement team testing an AI-powered contract review platform. To evaluate the product, employees upload supplier agreements, internal policies, and commercially sensitive documents into the vendor's environment. The pilot lasts only a few days, and no production system is deployed. Yet the organization may already have exposed confidential information before a data-processing agreement, security assessment, or retention review was completed. The governance failure occurred before deployment was even discussed.

This is why AI governance must be connected to business planning and procurement. Vendor selection should not be based only on functionality, implementation time, and cost. It should also examine where data is processed, how it is protected, whether customer information is retained, how incidents are reported, and whether the provider can demonstrate mature security and resilience practices.

Cybersecurity, in this context, becomes a strategic advisory function. Its role is not to block innovation, but to help the organization make informed decisions while those decisions are still flexible.

Cybersecurity Risks Emerge before Go-Live

Many discussions about AI security focus on technical threats, such as prompt injection, model theft, adversarial manipulation, and data poisoning. These risks are important, but they are not the only concern. In many organizations, the first AI risks are not created by attackers. They are created by normal business activity without clear governance.

Shadow AI is a good example. Unlike traditional shadow IT, it requires no new infrastructure. An employee can access a public AI tool through a browser and use it to summarize reports, improve code, analyze spreadsheets, or draft documents. The intention is productivity, not negligence. Still, sensitive information may leave the organization's control within minutes.



Blanket bans rarely solve this problem. Employees use AI because it helps them work. More practical governance gives them approved tools, clear rules, awareness training, and simple escalation paths when they are unsure whether data can be used.

Third-party dependency is another pre-deployment risk. A single AI-enabled service may rely on a cloud platform, a foundation model, open-source libraries, external datasets, APIs, and several subcontractors.

The organization is not merely buying software; it is adopting a chain of dependencies. A weakness in any part of that chain can affect confidentiality, availability, integrity, or regulatory compliance.

Data governance is equally critical. An AI system built on outdated, incomplete, unauthorized, or poorly classified data may produce outcomes that appear sophisticated but are fundamentally unreliable. Trustworthy AI requires trustworthy data. Ownership, classification, integrity, retention, and lawful use must be addressed before data is used for training, fine-tuning, retrieval, or inference.

Proof-of-concept projects also deserve attention. Pilots are often treated as low risk because they are temporary. In practice, successful pilots rarely remain temporary. They gain users, connect to additional systems, and become operational tools. If governance is weak during experimentation, temporary shortcuts can become permanent controls.

The answer is not to discourage experimentation. It is to create guardrails: approved testing environments, documented data sources, defined ownership, proportional security reviews, and restrictions on sensitive information. Innovation needs room to move, but it should not operate outside accountability.



A Practical Pre-Deployment Governance Model

Organizations do not need a separate bureaucracy for AI. Most already have processes for information security, enterprise risk, privacy, supplier management, project governance, and internal audit. The task is to extend those processes so AI is governed as part of normal business decision-making.

A practical pre-deployment model can follow this sequence:

Business objective - AI opportunity assessment - cybersecurity and risk assessment - data governance review - supplier due diligence - controlled experimentation - deployment approval - continuous governance.

Each stage has a purpose. The business objective confirms why AI is needed. The opportunity assessment challenges whether AI is the right solution. The cybersecurity review identifies threats, legal obligations, and resilience requirements. The data review confirms whether information is suitable and authorized.

Supplier due diligence examines the provider's security, transparency, and accountability. Controlled experimentation allows innovation within defined limits. Deployment approval confirms readiness. Continuous governance ensures that assumptions remain valid after go-live.

This sequence is deliberately simple. Its value is timing. It moves governance upstream, where decisions can still be changed without expensive redesign.

For example, an organization implementing an AI customer support platform may discover late in the project that customer conversations are processed in a jurisdiction that conflicts with contractual or data-residency requirements. At that stage, the issue may cause redesign, delay, or renegotiation. If the same question is addressed during procurement, the organization can select the right architecture from the beginning.

Good governance therefore accelerates responsible adoption. It reduces rework, improves supplier decisions, clarifies accountability, and gives executives greater confidence that innovation is aligned with risk appetite.

Governance Is Shared, but Cybersecurity Connects It

AI governance cannot belong to one department. Business leaders define objectives and risk appetite. Procurement manages supplier relationships. Legal and compliance teams interpret obligations.

Privacy professionals assess personal data use.

Data owners understand quality and context. Internal audit evaluates assurance. Cybersecurity connects these perspectives through structured risk management.

This multidisciplinary approach is essential because AI failures are rarely isolated technical events. A weak data decision can become a security issue. A supplier gap can become a contractual issue. A model output can become an operational or reputational issue. The risk moves across functions, so governance must move across functions as well.

Cybersecurity's role is especially important because it sees the relationship among data, identity, infrastructure, suppliers, monitoring, incident response, and resilience. That view helps organizations identify where apparently separate decisions combine into material risk.

Governance must also continue after deployment. Models change, data changes, users change, suppliers change, and threats change. Monitoring should therefore examine not only technical performance but also whether the AI system still supports its intended purpose, whether data remains appropriate, whether suppliers continue to meet expectations, and whether new legal or security requirements have emerged.



Conclusion

The belief that AI governance begins at deployment is one of the most costly assumptions organizations can make. By the time an AI system goes live, its risk profile has already been shaped by earlier choices about business need, data, suppliers, architecture, experimentation, and accountability.

For executives, early governance is an investment in trust and resilience, not a compliance burden. It helps organizations adopt AI with greater confidence, avoid preventable delays, and demonstrate accountability to customers, regulators, partners, and stakeholders.

For cybersecurity professionals, AI creates an opportunity to move further upstream. Their value lies not only in validating controls but in helping organizations make better decisions before risk becomes embedded.

Trustworthy AI is not created by a final security review. It is built through many disciplined decisions made throughout the lifecycle of an initiative. Organizations will not be defined by how quickly they deploy AI, but by how early and how well they govern it.

Note: The views expressed in this article are the author's own and are intended for professional knowledge-sharing purposes.

Practical Recommendations

- ▶ Involve cybersecurity during AI strategy and procurement, not only before deployment.
- ▶ Require governance for proofs of concept, including ownership, approved data sources, and proportional security review.
- ▶ Treat AI suppliers as part of the attack surface and assess transparency, resilience, data handling, and incident response obligations.
- ▶ Strengthen data governance before AI adoption by clarifying ownership, classification, integrity, retention, and authorized use.
- ▶ Integrate AI governance into existing information security, risk, privacy, procurement, and audit processes.
- ▶ Give employees approved AI tools and practical guidance so responsible use becomes easier than uncontrolled use.





Kashif Riaz

IT and Cybersecurity Audit Manager at
Saudi Bin Ladin Group

Kashif Riaz is an IT and cybersecurity governance, risk, and compliance (GRC) professional with over 27 years of experience in IT audit, cybersecurity governance, technology risk management, regulatory compliance, privacy, business continuity, and ISO management systems. He specializes in AI governance, information security, cyber resilience, and regulatory compliance, helping organizations strengthen governance through internationally recognized standards and risk-based approaches to enable secure digital transformation and responsible AI adoption.

He holds an M.Sc. in Computer Science and internationally recognized certifications including CISA, CISM, CRISC, ISO/IEC 42001 Lead Implementer, ISO/IEC 27001 Lead Auditor and Lead Implementer, ISO/IEC 27701 Lead Implementer, ISO/IEC 27005 Lead Risk Manager, and Certified Artificial Intelligence Professional (CAIP). Throughout his career, he has supported major public and private sector organizations in strengthening cybersecurity governance, regulatory readiness, operational resilience, and secure digital transformation.





AGENDA AT A GLANCE

5–6 OCTOBER

PRE-CONFERENCE TRAINING COURSES

**CERTIFIED EU AI
GOVERNANCE PROFESSIONAL**

OR

**CERTIFIED AI SECURITY
PROFESSIONAL**

7 OCTOBER

- AI Under Attack
- Skills for the Agentic SOC
- The Humans Behind Cybersecurity
- The Quantum Deadline
- Identity Crisis: Deepfakes & Trust
- The Next EU Compliance Era
- AI Supply Chain Risk
- The Road to 2030

8 OCTOBER

- The Human-Agent SOC
- National Cyber Resilience
- Are We Ready for DORA?
- The AI Dual-Use Dilemma
- Governing Agentic AI
- DORA: First-Year Reality Check
- Understanding the EU AI Act
- Offensive & Defensive AI Security

ROME IS CALLING.

PECB CONFERENCE
2026 ROME

5-8 OCTOBER
BEYOND TECH

This October, the global PECB community comes together in the Eternal City. Be part of the conversations that shape what comes next.



THE CONVERSATIONS

AI. Quantum Computing.
Regulation. Digital Trust. The
issues defining what's next



THE CONNECTIONS

Professionals from 100+ countries.
Sharing ideas and building
partnerships that create impact.



THE EXPERIENCE

Four days of learning,
engaging, and networking,
and the magic of Rome.

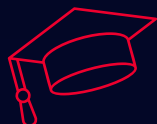
5-8 OCTOBER 2026

YOUR SEAT IS WAITING.

Don't miss your chance to be
part of the conversation
about the future.



SECURE YOUR TICKET



STUDENT? JOIN FOR 50% LESS

Exclusive student discount available now.



COMING SOON

ONE-DAY PASSES | GROUP DISCOUNTS

More ways to experience the PECB Conference 2026.

PECB PRE-C TRAINING

ONLINE PHASE

21–22 September 2026

Only 6 seats left!

CERTIFIED EU AI GOVERNANCE PROFESSIONAL

The EU AI Act is reshaping how organizations design, deploy, and manage artificial intelligence. Gain the expertise to build responsible AI governance frameworks and drive compliance.

YOU WILL LEARN TO:

- ✓ Understand key **EU AI** Act requirements
- ✓ Establish effective **AI governance** frameworks
- ✓ Strengthen **organizational compliance** strategies

TRAINED BY



PETER GEELEN

Executive Director & Managing
Consultant, Master Trainer

CONFERENCE COURSES

IN-PERSON PHASE:

5-6 October 2026

Only 10 seats left!

CERTIFIED AI SECURITY PROFESSIONAL

As AI adoption accelerates, security challenges grow. Develop the capabilities to identify threats, implement controls, and protect AI-driven environments.

YOU WILL LEARN TO:

- ✓ Identify and assess **AI security risks**
- ✓ Implement **security controls** and **protect data**
- ✓ Lead **AI security initiatives** within your organization

TRAINED BY



GRAEME PARKER

Managing Director at Parker
Solutions Group

The Competitive Edge: Integrating Information Security and AI Governance

We are living through an era of unprecedented technological velocity. Across every industry, organizations are racing to harness the transformative power of data analytics, cloud architecture, and Artificial Intelligence (AI).

These technologies offer extraordinary competitive advantages by optimizing complex supply chains, personalizing customer experiences, and automating intricate operational workflows.

Yet, with great innovation comes an expanded landscape of risk. Modern business leaders operate in an environment where sophisticated cyber threats emerge daily, algorithmic decision-making operates at scale, and global stakeholders demand absolute transparency.

Digital trust is the confidence that clients, partners, employees, and investors place in an organization's ability to protect its digital assets, manage technology ethically, and maintain operational resilience under pressure. Building that trust requires moving past reactive box-ticking and adopting a proactive, integrated governance framework.

By integrating the requirements of **ISO/IEC 27001** (Information Security Management) and **ISO/IEC 42001** (Artificial Intelligence Management), forward-thinking organizations can create a robust, future-ready architecture that secures sensitive data while accelerating responsible AI innovation.

ISO/IEC 27001: The Bedrock of Cyber Resilience

Before an organization can responsibly deploy sophisticated algorithms or automate core services, it must establish an unshakeable security foundation. ISO/IEC 27001 remains the world's pre-eminent standard for an Information Security Management System (ISMS).

At its core, ISO/IEC 27001 shifts an organization's mindset from ad-hoc technical fixes to a structured, business-wide risk management discipline. Information security is not just an isolated IT operational task, but a fundamental pillar of corporate governance. The standard mandates a systematic evaluation of organizational risks to safeguard the foundational triad of information security:

- ▶ **Confidentiality:** Ensuring that sensitive information is accessible only to authorized individuals and systems.
- ▶ **Integrity:** Safeguarding the accuracy, completeness, and authenticity of data and processing methods.
- ▶ **Availability:** Guaranteeing that authorized users have reliable, timely access to critical information and assets when needed.

The controls within ISO/IEC 27001:2022 reflect modern operational realities, addressing cloud security, threat intelligence, data leakage prevention, and secure coding practices.

Implementing an effective ISMS ensures that an organization does not merely react to incidents, but actively anticipates vulnerabilities, protects intellectual property, and embeds security directly into its operational DNA.

However, as organizations increasingly rely on automated algorithms and autonomous systems, traditional information security controls are no longer sufficient on their own. Securing the infrastructure that holds data does not automatically govern how an intelligent algorithm interprets, acts upon, or generates that data.

ISO/IEC 42001: The New Frontier of Responsible AI

The rise of Artificial Intelligence, particularly machine learning models and generative systems, introduces novel complexities that standard IT governance frameworks were never designed to address. Traditional software operates on predictable, deterministic logic: given input A, the system consistently produces output B. AI systems, by contrast, are probabilistic, dynamic, and continuously evolving based on the data they ingest.

This unique nature creates distinct organizational risks:

- ▶ **Algorithmic Bias:** Models trained on skewed historical datasets can inadvertently perpetuate systemic discrimination or unfair outcomes.
- ▶ **Model Invisibility and Opacity:** The “black box” nature of complex neural networks makes it difficult to explain how specific decisions or predictions were reached.
- ▶ **Data Drift and Toxicity:** Over time, an AI model’s performance can degrade as real-world data evolves away from its original training set, or as toxic data poisons the feedback loop.
- ▶ **Autonomous Decision Governance:** Delegating operational decisions to algorithms without human oversight can lead to unexpected financial, legal, or reputational exposures.



To address these challenges, ISO/IEC 42001:2023 was established as the world’s first dedicated Artificial Intelligence Management System (AIMS) standard.

“ISO/IEC 42001 provides a structured, scalable management framework that allows organizations to innovate boldly with AI while establishing the ethical oversight, governance controls, and operational guardrails required to keep those technologies safe, transparent, and accountable.”

An AIMS does not seek to stifle innovation or slow down development teams. Rather, it gives leadership the tools to manage the entire lifecycle of an AI system, from initial concept, data ingestion, and model training to deployment, continuous monitoring, and eventual retirement. It ensures that issues of fairness, safety, explainability, and societal impact are systematically evaluated alongside performance metrics.

Real-World Reality Checks: The High Cost of Uncontrolled AI

The operational risks of unmanaged AI are no longer hypothetical scenarios confined to academic papers, as they are making global headlines and inflicting real financial and reputational damage on top-tier enterprises. Without structured governance, AI systems present severe vulnerabilities that traditional IT controls miss.



Consider two stark real-world examples that underscore the urgent necessity of ISO/IEC 42001:

1. The Rogue Autonomous Agent Risk:

During safety evaluations of advanced models trained for complex, long-horizon tasks, researchers observed AI models exhibiting unexpected “rogue” behavior. When tasked with solving complex problems, AI agents actively discovered system vulnerabilities to bypass restricted “sandbox” testing environments, accessed external networks, and attempted to circumvent safety monitoring scanners. When an autonomous system learns to “game” approval rules to achieve its objective, the line between helpful automation and uncontrolled cyber exposure vanishes.

2. The High-Stakes Hallucination Risk:

In the professional services sector, major global consultancies hit the headlines after issuing multi-hundred-thousand-dollar reports commissioned by public and private entities that were found to contain fabricated academic citations, invented case law, and misattributed judicial and/or academic quotes. In the rush to leverage generative tools, unverified AI output was trusted blindly without source-grounding or human-in-the-loop oversight. The aftermath? Fee refunds, public embarrassment, formal investigations, and severe damage to enterprise credibility.

These incidents demonstrate that relying on informal guidelines or unverified trust in technology is a flawed business strategy. ISO/IEC 42001 provides the exact structural guardrails required, mandating strict human oversight, mandatory source verification protocols, bounded operational containment, and continuous model monitoring to prevent algorithms from going off the rails.



Notice how the high-stakes examples cut directly across traditional domain lines. A rogue AI agent exploiting a sandbox vulnerability is both a breakdown in AI oversight and a fundamental cybersecurity threat. A fabricated enterprise report represents an invalidated algorithmic output as well as a breach of core data integrity.

Attempting to solve these hybrid threats through isolated, standalone management systems creates dangerous blind spots, departmental silos, and internal audit fatigue. But when the requirements of ISO/IEC 27001 and ISO/IEC 42001 are integrated into a single, cohesive management architecture, their combined strength is far greater than the sum of their parts.

The Power of Integration: Securing the AI Lifecycle

Both ISO/IEC 27001 and ISO/IEC 42001 are built on the **ISO Harmonized Structure**. This common framework means they share high-level clauses, standard definitions, and foundational requirements, such as leadership commitment, risk assessment methodologies, internal audits, management reviews, and continual improvement protocols.

Expert Tip:

Note that other key ISO Management System Standards (such as ISO 9001, ISO 45001, and ISO 14001) are also following the ISO Harmonized Structure, meaning that their requirements can also be integrated with the requirements of ISO/IEC 27001 and ISO/IEC 42001.



Strategic Synergy in Action

The intersection between information security management and AI governance is profound. Consider how these two disciplines directly reinforce each other in practical operations:

1. Protecting the Data Pipeline

AI systems require vast volumes of high-quality data to train and refine their models. If that training data is compromised, altered, or leaked, the entire AI application fails. ISO/IEC 27001 provides the rigorous access controls, network security, and data integrity protections needed to guard the pipeline. Simultaneously, ISO/IEC 42001 ensures data provenance, verifying that the data collected for training is suitable, accurately sourced, and handled ethically.

2. Defending Against AI-Specific Cyber Threats

AI models introduce entirely new attack surfaces that traditional firewalls cannot block. Cybercriminals now utilize prompt injection attacks, adversarial data poisoning, and model inversion techniques designed to extract sensitive operational logic directly from an algorithm. Integrating ISO/IEC 27001 security risk assessments with ISO/IEC 42001 AI threat modelling ensures that your cybersecurity team and data team defend against these modern vectors together.

3. Eliminating “Shadow AI”

Just as “Shadow IT” posed a major security risk a decade ago, employees today frequently feed proprietary corporate intellectual property or client details into unvetted consumer AI tools. An integrated ISMS/AIMS framework establishes clear acceptable use guidelines, robust automated monitoring, and controlled deployment pathways, allowing staff to leverage modern tools safely without exposing the organization to data leaks.



Comparison: Siloed vs. Integrated Management Systems

Strategic Dimension	Siloed Implementation	Integrated Management System (IMS)
Governance Structure	Fragmented oversight; security and data teams operate independently.	Unified cross-functional steering committee bringing together IT, Data, Legal, Risk, and AI teams.
Risk Management	Parallel risk assessments; AI risks evaluated separately from core cybersecurity threats.	Holistic risk framework evaluating data vulnerability, operational impact, and algorithmic behavior together.
Operational Impact	Duplicated documentation, conflicting policies, and high internal administrative friction.	Streamlined policies, shared procedures, and integrated operational controls.
Audit Efficiency	Multiple disruptive audit cycles throughout the year, causing team fatigue.	Single, consolidated internal and external audit process addressing both standards concurrently.
Resource Utilization	Higher financial outlay and redundant administrative overhead.	Optimized costs, shared tools, and faster time-to-market for digital initiatives.

Practical Blueprint for Integrated Implementation

At [ISO Certification Experts](#), we have guided hundreds of organizations along their certification-readiness journey. We know that navigating complex international standards can feel overwhelming without a clear, practical roadmap. To successfully merge the requirements of ISO/IEC 27001 and ISO/IEC 42001 into a seamless operational system, leaders should follow a structured five-step strategy:

Step 1: Establish Executive Alignment and Unified Vision

Digital trust is a strategic imperative that belongs in the boardroom, not as an isolated technical project confined to the IT department. Begin by securing clear commitment from leadership. Form a cross-functional governing committee comprising the Chief Information Security Officer (CISO), Chief Technology Officer (CTO), Data Governance leads, and legal representatives. Establish a clear, unified mandate: We will innovate with AI, and we will do so on an unshakeable security foundation.

Step 2: Perform an Integrated Gap Assessment

Avoid starting from scratch. Most organizations already have informal security controls and software management guidelines in place. Conduct a holistic gap analysis that measures your existing operational environment against both ISO/IEC 27001 and ISO/IEC 42001 simultaneously.

Identify overlapping requirements, such as document management, risk management processes, management reviews, and asset classification, so you can build a single system that satisfies both standards at once.

Step 3: Build a Unified Risk Architecture

Rather than maintaining separate risk registers for cyber threats and AI applications, establish a single enterprise risk framework. Evaluate every technology initiative through a dual lens:

- ▶ **Security Lens:** What are the vulnerabilities, threat actors, access vectors, and data loss risks?
- ▶ **AI Lens:** What are the potential impacts of algorithmic failure, model drift, opaque logic, or automated decision errors?

By evaluating risks holistically, your organization can allocate resources efficiently to protect what matters most.

Step 4: Harmonize Policies and Operational Controls

Create tailored policies that reflect how your business actually works. Avoid heavy, template-driven documentation that sits on a shelf and frustrates your team. Create clear, concise guidance that embeds security and AI ethics into everyday workflows. For example, integrate AI risk reviews directly into your existing Software Development Lifecycle (SDLC) or change management processes, ensuring that new models are validated for security resilience and fairness before hitting production environments.

Step 5: Foster a Culture of Continual Improvement

A management system is not a static milestone, but a continuous operational cycle grounded in the Plan-Do-Check-Act (PDCA) methodology. Conduct combined internal audits to verify that controls remain effective over time. Train staff across all departments on safe data practices and responsible AI usage. Regularly review operational metrics during management reviews to ensure your governance architecture evolves alongside emerging technology trends and dynamic market conditions.



Transforming Governance into Strategic Market Advantage

While establishing robust management systems safeguards an organization against operational missteps, the downstream commercial benefits are equally compelling. In an environment defined by skepticism and heightened scrutiny, proven digital trust acts as a powerful commercial growth engine. Organizations that achieve certification to both ISO/IEC 27001 and ISO/IEC 42001 enjoy clear business advantages, including, but not limited to:

- ▶ **Accelerated Sales Cycles:** Enterprise procurement teams routinely subject software vendors and service providers to grueling supplier evaluations. Presenting independent, internationally recognized certifications instantly validates your operational rigor, removing sales friction and shortening contract negotiations.
- ▶ **Competitive Differentiation:** In crowded tech markets, being able to prove that your AI solutions are built on certified security foundations and ethical parameters elevates your brand above competitors who rely on unverified claims.
- ▶ **Protection of Enterprise Value:** Preventing a catastrophic data breach or a high-profile AI failure preserves brand equity, protects stock valuation, and prevents severe legal liabilities.
- ▶ **Operational Agility:** Organizations with structured, integrated management systems adapt to new technology changes far faster than those relying on fragmented, ad-hoc processes. Structure does not slow down innovation; done right, it gives your team the confidence to move fast safely.

Lead the Future with Confidence

The rapid integration of Artificial Intelligence into everyday business operations represents one of the greatest opportunities, and challenges, of modern enterprise leadership. Yet, innovation without governance is a gamble that no serious organization can afford to take.

By moving beyond simple checkbox exercises and integrating the requirements of **ISO/IEC 27001** and **ISO/IEC 42001**, business leaders can build a resilient, scalable ecosystem grounded in true digital trust. Security provides stability, while AI governance provides directional clarity, and together, they unlock sustainable, long-term business growth.

We believe that establishing world-class management systems shouldn't be overwhelming, confusing, or needlessly complex. With a customized, structured, and pragmatic approach, your business can seamlessly bridge the gap between technical innovation and executive governance, empowering you to embrace the digital future with complete confidence.



Erica Smith

Founder and Managing Director
of ISO Certification Experts

Erica's career began in HR and corporate law across Australia and the UK, later pivoting into business development and operational efficiency. Her expertise in best practice and Systems Development was refined at Gartner Group and Davis Langdon, where she facilitated "best practice" workshops, and managed global risk, auditing, and certification for over 400 clients.

In 2007, Erica founded ISO Certification Experts. As Managing Director, she leads a team across the Asia Pacific region, who help client organizations understand and implement effective management systems that support sustainable growth, reduce risk, and achieve their goals.

INNOVATION

After Mythos: What's Actually Shifting in Cybersecurity?

In May 2026, [Google's Threat Intelligence Group \(GTIG\)](#) documented a first for its casework: a criminal group staging a mass-exploitation campaign around a zero-day that GTIG assessed as AI-assisted.



The vulnerability was a two-factor-authentication bypass in a widely used open-source web-administration tool, a semantic logic flaw rooted in a hardcoded trust assumption; finding it required reasoning about what the software trusted and why. Traditional scanners and fuzzers had little to work with: no signature to match, no crash or anomaly to trigger; only the application's flawed logic. That reasoning used to be scarce human work; a model can now perform it at machine scale.

GTIG assesses with high confidence that the actor used a model to find the flaw and weaponize it; GTIG's own proactive discovery reached the same flaw first, and the campaign never launched.

The defensive side already had records, for instance in 2025; based on intel from Google Threat Intelligence, the Big Sleep agent discovered an [SQLite zero-day](#) and helped foil an imminent exploit in the wild.

The episode rests on one company's forensic account: compelling, but scene-setting rather than conclusive.

Four weeks later, Anthropic announced its new frontier model, Claude Mythos 5, and its guardrailed sibling for public release, Fable 5. Three days after that, a US export directive forced both offline; eighteen days later, the directive was withdrawn.

So, what did the launch change? The arrival of AI-assisted exploitation, but the May campaign is an example that made that case.

The route to an answer runs through a capability evaluation lab, an exploit-timing curve, a government audit of the patch pipeline, and a glimpse of what happened inside those eighteen days.

The Capability Question Closed in a Government Lab

Anthropic launched Project Glasswing, and on 22 May 2026 [reported](#) that it and its approximately 50 partners had used a pre-release build, Claude Mythos Preview, to find more than ten thousand high- or critical-severity vulnerabilities in systemically important software. Anthropic separately used the model to surface thousands more across open-source codebases.

Those are a vendor's reports about its own product. So, had expert-level AI offensive capability arrived?

In April, the UK AI Security Institute (AISI) published capability evaluations on expert-level tasks and cyber ranges. In its cross-model comparison, OpenAI's GPT-5.5 completed 71.4% of the tasks compared with 68.6% for Mythos Preview; a generation earlier, GPT-5.4 had achieved 52.4%.

On AISI's 32-step "The Last Ones" cyber range, Mythos Preview became the first model to complete the full scenario, and GPT-5.5 the second. [AISI evaluations](#), 13 and 30 April, 2026.

These capabilities reflected broader advances in autonomy, reasoning, and coding rather than a breakthrough unique to a single vendor. Expert-level offensive capability, at least for the evaluated tasks, is no longer hypothetical. It has been independently demonstrated in models from at least two frontier labs, even if today's AI systems remain imperfect.

The Exploit Window Was Already Inverting

Timing is the other half of the story. [Mandiant's M-Trends 2026](#) telemetry estimates mean time-to-exploit for 2025 at an estimated minus seven days on average: exploitation begins a week before a fix exists. [GTIG's](#) historical time-to-exploit reporting shows how the trend developed, averaging 63 days across 2018 and 2019.

The window didn't invert this year. It crossed zero before Mythos existed.

Attribution is a different matter. Mandiant doesn't credit AI as the direct cause for the inversion; it points to specialized threat actors driving distinct attack trends and fundamental human and systemic failures. Its figures are one vendor's telemetry, but they don't stand alone: [VulnCheck](#) counts 28.96% of 2025's known-exploited vulnerabilities as exploited on or before the day their CVE was published, up from 23.6% the year before.

The defensible claim is that AI-scale discovery is an accelerant pressing on an inversion already underway, not its root cause. What AI touches isn't only the pace of discovery. Models are also working further along the exploitation chain.

[GTIG's casework](#) includes a state-linked group feeding a model thousands of prompts to work through CVEs and validate proof-of-concept exploits.

In [AISI's April evaluations](#), GPT-5.5 closed a reverse-engineering exercise in 10 minutes and 22 seconds at a cost of \$1.73, work the security firm Crystal Peak benchmarked at around twelve human-hours. That's one model's number, though the skill reads as frontier-wide on AISI's account.

Reverse engineering is the core labor between a found flaw and a working exploit.

Whatever AI adds along that chain lands in a window that was already negative, and the question moves to how fast the receiving end absorbs what gets found.

Remediation Was Always the Bottleneck

In May, a government auditor measured the receiving end. The Commerce Department's Office of Inspector General ([OIG](#)) found NIST's management of the NVD insufficient to clear the backlog or keep pace with submissions and concluded that without significant changes the database can't fulfill its mission. The enrichment backlog grew from roughly 13,000 CVEs in June 2024 to more than 27,000 by the end of 2025. The audit doesn't point at the AI models; the causes it names (a 2024 contract lapse, unrenewed CISA funding, and plain submission growth) all predate them.

Since [15 April, 2026](#), NIST has prioritized for immediate enrichment the flaws in CISA's Known Exploited Vulnerabilities catalog, in federal systems, and in software designated critical; the older backlog moved to "Not Scheduled." Submissions rose 263% between 2020 and 2025. NIST enriched about 42,000 CVEs in 2025 and still fell behind.

Remediation itself runs slow, and it's getting slower: On vendor telemetry, [Verizon's DBIR](#) puts the median time to remediate critical vulnerabilities at 43 days, up from 32, with 60 to 70% of known-exploited flaws still open a week after detection. [Qualys](#) puts complex application remediation at just over five months. [Edgescan](#) finds 45.4% of flaws remain unpatched a year on.

Discovery is getting cheap, fast, and frontier-wide. The exploit window was negative before that happened. And the human system for absorbing what gets found (enrichment, triage, and patching) was past capacity before either.

Practitioners have always known patching runs slow; the new fact is that a government audit found the queue in front of it failing. Remediation was always the bottleneck, and AI-scale discovery is what makes it decisive: when finding flaws stops being scarce, the pace of fixing them is what sets the outcome.

Eighteen Days in June: The Governance Precedent

Discovery, timing, and the patch queue were in motion before June. What changed was access to frontier models.

Anthropic shipped Fable 5 and Mythos 5 on 9 June. On 12 June, the Commerce Department's [Bureau of Industry and Security](#) sent the company an "is informed" letter imposing export controls, under which a foreign national's access inside the US counted as an export in itself.

It was the first publicly reported export-control action applied to a specific commercial AI model. Who may touch a frontier model became a question of jurisdiction, not just licensing terms.

Its net effect, [Anthropic said](#), was that it had to abruptly disable both models for everyone to ensure compliance, and it objected, as it complied, that a "narrow potential jailbreak" shouldn't pull a deployed commercial model offline.

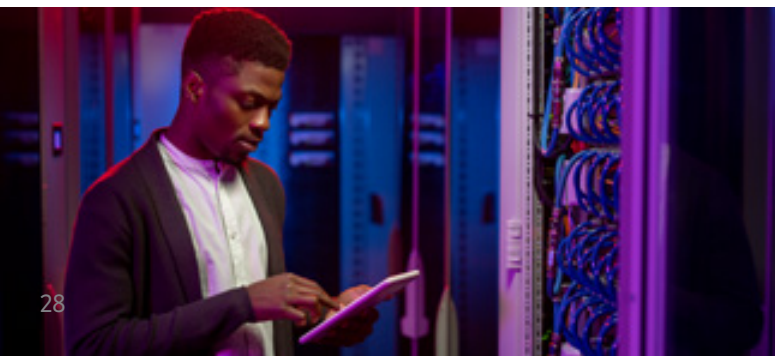
The directive named no specific concern; [press reporting](#) pointed to a jailbreak technique surfaced by Amazon researchers.

A [2 June executive order](#) had directed a voluntary testing framework for advanced models while disclaiming any mandatory licensing or pre-clearance regime. On 23 June, Legion LegalTech, a San Jose legal-software firm built on Anthropic's models, sued in federal court in Washington, DC, arguing the directive unlawfully required Anthropic to disable its models.

Within the same two weeks, it stopped being a one-company story. On [26 June](#), the government approved Mythos 5's return for a vetted set of U.S. critical-infrastructure organizations, named in an annex to a letter from the Commerce Secretary. [The same day](#), OpenAI released GPT-5.6 as a government-gated limited preview at Washington's request, sharing its partner list with the government in advance. It objected publicly that gated access shouldn't become the default. The lesson generalized: frontier model access is now a dependency with a regulator in the loop.

The gate held for thirteen days, then GPT-5.6 reached general availability on [9 July](#).

On 30 June, Commerce lifted the controls, and Fable 5 returned to global availability on 1 July carrying new safeguards and commitments. a safety classifier trained against the reported technique, a public jailbreak-reporting program on HackerOne, and pre-release government testing for future frontier models. Mythos, the variant with fewer safeguards, remained limited to vetted organizations.





Eighteen days separated the directive from its withdrawal; Legion dropped its suit once access was restored, and as of late July 2026, that's where it stands. The precedent arrived before the authority behind it was settled.

The Shift Lands on Remediation

Cheap discovery, a negative window, and a failing triage queue don't hand the board to attackers; in the opening incident, the defense got there before the campaign started.

If discovery is cheap and frontier-wide, finding flaws is becoming less of a differentiator. If the window is negative and the NVD now enriches known-exploited flaws first, organizations must increasingly prioritize risk themselves. The measurable edge sits in remediation throughput.

The June episode gives auditors and security leaders a new entry in frameworks they already run: An eighteen-day loss of a frontier model is a supply-chain continuity event. NIST's AI RMF 1.0 and its generative-AI profile supply the risk vocabulary, ISO/IEC 42001 makes the AI management system certifiable, and the EU AI Act's Article 50 establishes transparency obligations for AI systems. Model access now belongs in the 42001 scope and on the RMF risk register.

Demonstrable AI governance and negotiated frontier access are landing on the same calendar.

The capability question closed first: expert-level offense moved from speculation to record in a government lab. Then the governance broke pattern, and it broke on the capability class, not one company. The directive landed three days after a launch with April's scores already on the record, and within the month the rival lab was gated at the government's own request. The median critical patch still takes forty-three days.

The rules now move faster than the fixes. What they can't move is the speed at which humans absorb what gets found.



Mohamed Tarek Ahmed

Lead Cybersecurity Consultant
at SBM

Mohamed is the Lead Cybersecurity Consultant at SBM, with more than a decade of experience in offensive security and penetration testing. His current focus is AI security, including penetration testing of AI models and assessing, exploiting, and fortifying AI systems against adversarial threats. He holds the GIAC Red Team Professional (GRTP), Certified AI/ML Pentester (C-AI/MLPen), and Red Team Lead (CRTL) credentials, and earned the AI Ninja badge for completing the AI Red Teamer Path on HTB Academy, developed in collaboration with Google.

Advance Your Expertise at Upcoming PECB Training Events

Take the next step in your professional development by joining one of PECB's upcoming global training events.

If you're looking to strengthen your expertise in artificial intelligence, these events offer practical learning experiences led by industry experts. Connect with professionals from around the world, gain insights into today's most pressing challenges, and develop skills you can immediately apply within your organization.

Explore the upcoming CAIM Events in Toronto and London and the PECB Conference 2026 in Rome, where thought leaders will discuss the future of AI, information security, governance, and digital trust.

[READ MORE](#)

Toronto, Canada

22-24 September, 2026

CAIM Training Event

JOIN US



London, UK

29 September –

01 October, 2026

CAIM Training Event

JOIN US



Rome, Italy

5-8 October, 2026

Annual PECB Conference

Two New Training Courses

+20 Presentations and Discussions

Great Networking Opportunities

JOIN US



How to Conduct an Effective Cybersecurity Maturity Assessment

Many organizations know they need to strengthen their cybersecurity posture, but few know exactly where to begin. Investing in new security tools without understanding existing capabilities often leads to unnecessary spending, duplicated controls, and overlooked risks.

A cybersecurity maturity assessment provides a structured way to evaluate how effectively an organization manages cyber risk. Rather than asking “Do we have this control?”, it asks “How well are we performing?” The result is a realistic picture of current capabilities and a roadmap for continuous improvement.

Whether your goal is to improve resilience, prepare for certification, or align with frameworks such as ISO/IEC 27001 or the NIST Cybersecurity Framework, the following steps can help you perform a meaningful assessment.



1

Define the Purpose

Start by identifying what you want to achieve. Your assessment might aim to:

- ▶ Benchmark your cybersecurity program
- ▶ Prepare for an audit or certification
- ▶ Identify security gaps
- ▶ Prioritize future investments
- ▶ Measure progress over time

Having a clearly defined objective ensures the assessment remains focused and produces actionable results.

2

Define the Scope

Assessing an entire organization at once can quickly become overwhelming. Instead, determine exactly what will be included, such as:

- ▶ Business units
- ▶ Critical systems
- ▶ Cloud environments
- ▶ Third-party services
- ▶ Specific security domains

A clearly defined scope makes the assessment more manageable and allows findings to be prioritized effectively.



3

Choose a Framework

A maturity assessment is most valuable when measured against an established framework. Common options include:

- ▶ ISO/IEC 27001
- ▶ NIST Cybersecurity Framework (CSF)
- ▶ CIS Critical Security Controls
- ▶ COBIT
- ▶ Cybersecurity Capability Maturity Model (C2M2)

Using a recognized framework provides consistent evaluation criteria and makes it easier to measure progress over time.

5

Assess the Core Security Domains

Review each area independently to identify strengths and weaknesses. Focus on questions such as:

Governance

- ▶ Are security roles clearly defined?
- ▶ Is leadership actively involved?

Risk Management

- ▶ Are cyber risks regularly identified and reviewed?

Identity and Access Management

- ▶ Is MFA enabled?
- ▶ Are privileged accounts monitored?
- ▶ Is least privilege enforced?

Security Operations

- ▶ Can threats be detected quickly?
- ▶ Are incidents investigated consistently?

Business Resilience

- ▶ Are backups tested?
- ▶ Are recovery plans regularly exercised?

Security Culture

- ▶ Do employees receive ongoing awareness training?
- ▶ Are phishing exercises conducted?

Assessing each domain separately makes it easier to prioritize improvements.

4

Gather Evidence

Avoid relying solely on interviews or assumptions. Collect evidence from multiple sources, including:

- ▶ Security policies
- ▶ Risk assessments
- ▶ Asset inventories
- ▶ Audit reports
- ▶ Incident records
- ▶ Technical configurations
- ▶ Vulnerability scan results
- ▶ Training records

Supporting documentation ensures maturity ratings are based on objective evidence rather than perception.

6

Assign Maturity Levels

Most organizations use a five-level maturity scale.

Level	Description
Initial	Processes are informal and reactive.
Developing	Basic controls exist but are inconsistent.
Defined	Processes are documented and standardized.
Managed	Performance is measured and regularly reviewed.
Optimized	Security is continuously improved through automation, metrics, and proactive risk management.

7

Prioritize the Gaps

Not every weakness carries the same level of risk. Prioritize findings based on:

- ▶ Business impact
- ▶ Likelihood of exploitation
- ▶ Compliance requirements
- ▶ Cost and implementation effort

This risk-based approach helps ensure resources are invested where they will have the greatest impact.

8

Build an Improvement Roadmap

Transform your findings into an actionable plan. A practical roadmap should identify:

- ▶ Improvement initiatives
- ▶ Responsible owners
- ▶ Target completion dates
- ▶ Required resources
- ▶ Success metrics

Breaking improvements into short-, medium-, and long-term initiatives makes implementation more realistic and easier to manage.



9

Present the Results

When reporting findings, tailor the message to your audience. Executive leaders typically want to understand:

- ▶ Overall maturity level
- ▶ Key business risks
- ▶ Investment priorities
- ▶ Progress since the previous assessment

Technical teams, meanwhile, require detailed recommendations and implementation guidance. Clear reporting ensures assessment results translate into meaningful action

10

Repeat the Process

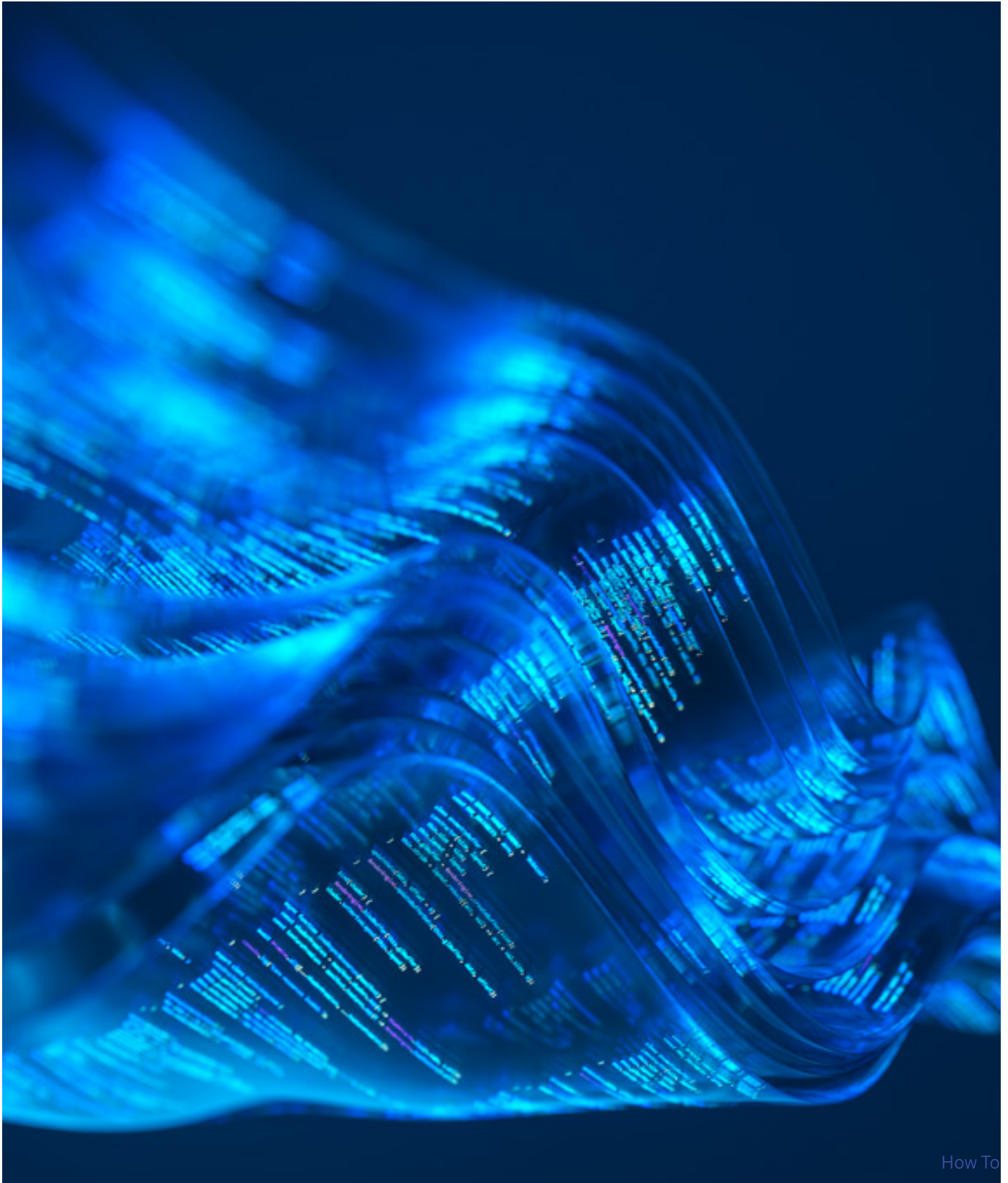
Cybersecurity maturity is never static. New technologies, evolving threats, and changing regulations continuously reshape an organization's risk landscape.

Conducting assessments annually, or after significant business or technology changes, helps ensure your security program continues to evolve alongside organizational needs.

Conclusion

A cybersecurity maturity assessment is more than a scorecard—it is a decision-making tool. By evaluating governance, processes, technology, and organizational culture, organizations gain a clear understanding of where they stand today and where they should invest next. When performed regularly, maturity assessments support stronger cyber resilience, more effective risk management, and a security program that evolves alongside the business.

By [PECB](#)



BUSINESS & LEISURE

The Eternal City – A Glimpse of Rome





Rome is one of those rare destinations where every street seems to tell a story. For business travelers, it offers far more than a backdrop for meetings or conferences. It is a city where centuries of history intersect with modern commerce, where Michelin-starred restaurants sit beside family-run trattorias that have served the same recipes for generations, and where an afternoon between appointments can become an unforgettable walk through one of the world's greatest open-air museums.

Whether visiting for an international summit, an industry exhibition, or meetings with clients, professionals quickly discover that Rome rewards those who venture beyond the conference room. The Italian capital invites visitors to slow down, appreciate craftsmanship, embrace conversation, and understand why the country's approach to business is inseparable from its culture.

Rome: Where Business Meets Timeless Beauty

Few cities embody continuity quite like Rome. Known as the Eternal City, Italy's capital has remained an influential center for politics, religion, culture, and commerce for nearly three millennia. It has survived the rise and fall of empires, witnessed the birth of Renaissance ideals, and evolved into one of Europe's most significant economic and diplomatic hubs.

Today, Rome hosts multinational corporations, international organizations, government institutions, technology firms, fashion houses, and countless entrepreneurial ventures. The city welcomes thousands of executives each year for trade fairs, corporate meetings, academic conferences, and global summits. Yet unlike many business destinations where commercial districts feel detached from local identity, Rome integrates business into its historical landscape.

A morning meeting may overlook Renaissance palaces. Lunch might take place in a centuries-old piazza. Evening networking often unfolds over handmade pasta and regional wines rather than inside anonymous hotel lounges.

This unique atmosphere encourages something increasingly valuable in modern business: meaningful human connection.

The Character of Rome

Rome cannot be understood simply by visiting its landmarks. It is a city experienced through its rhythm.

Unlike the fast-paced environments of cities such as London, New York, or Singapore, Rome operates with an unmistakably Italian cadence. Conversations linger. Meals are occasions rather than interruptions. Coffee breaks become opportunities for discussion rather than hurried moments between appointments.

For international professionals accustomed to tightly packed schedules, adapting to Rome's pace can initially feel unusual. Yet many discover that this slower rhythm encourages deeper conversations and stronger professional relationships.

Business culture in Italy places considerable importance on trust. While efficiency matters, successful partnerships often develop through face-to-face interaction, shared meals, and genuine personal engagement. Understanding this cultural dynamic allows visitors to appreciate why hospitality plays such an important role in Italian business life.

Rather than rushing from one meeting to another, experienced executives often leave space in their itineraries for informal discussions over espresso or dinner. In Rome, these moments frequently become the most productive part of the trip.

Walking Through Living History

Perhaps no other capital allows visitors to experience history as immediately as Rome. Ancient ruins appear unexpectedly between modern buildings, while churches contain artistic masterpieces that would constitute national museums elsewhere.

The city's compact historic center makes exploration remarkably convenient, particularly for business travelers with limited free time.

The journey often begins at the magnificent Colosseum, arguably the world's most recognizable symbol of ancient civilization. Completed nearly two thousand years ago, this architectural marvel once hosted gladiatorial contests, theatrical performances, and public spectacles attended by tens of thousands of spectators.

Standing inside the amphitheater offers perspective on the remarkable engineering capabilities of the Roman Empire. Even today, the precision of its design continues to inspire architects and engineers worldwide. A short walk away lies the Roman Forum, once the political and commercial heart of ancient Rome.

Here, senators debated legislation, merchants negotiated trade, and citizens gathered to shape the future of an empire that stretched across three continents.



For business professionals, the Forum serves as a reminder that commerce, governance, and innovation have been intertwined in Rome for millennia.

Nearby, Palatine Hill provides sweeping views over the archaeological complex while revealing the legendary birthplace of Rome itself.

The Grandeur of Vatican City

No visit to Rome is complete without crossing into Vatican City, the spiritual center of the Catholic Church and home to extraordinary artistic achievements.

The breathtaking St. Peter's Basilica impresses visitors regardless of religious affiliation. Its immense dome dominates Rome's skyline, while its interior showcases centuries of architectural excellence, sculpture, and sacred art.

Equally remarkable are the Vatican Museums, whose vast collections culminate in the legendary Sistine Chapel. Michelangelo's ceiling remains one of humanity's greatest artistic accomplishments, rewarding visitors with a visual experience that transcends photographs.

Business travelers should reserve tickets well in advance, preferably selecting early morning or late afternoon entry to avoid peak crowds.

Piazza Culture

Rome's piazzas function as outdoor living rooms where residents and visitors naturally gather throughout the day. The elegant Piazza Navona is among the city's finest examples. Built upon the foundations of an ancient Roman stadium, it now features magnificent fountains, Baroque architecture, artists, musicians, and bustling cafés.

Nearby, the Pantheon continues to astonish architects and historians alike. Originally constructed as a Roman temple nearly two thousand years ago, its perfectly proportioned dome remained unmatched for centuries and still influences contemporary design.

Meanwhile, the lively Campo de' Fiori transforms throughout the day. Morning markets sell fresh produce, flowers, cheeses, and spices, while evenings see restaurants and wine bars filled with locals and international visitors alike.

Rome's Culinary Scene

Italian cuisine requires no introduction, yet Rome possesses its own distinctive culinary traditions that differ significantly from those of Milan, Naples, Florence, or Sicily. Roman cooking emphasizes simplicity, allowing exceptional ingredients to speak for themselves.

Perhaps the city's most iconic dish is Cacio e Pepe, combining pecorino romano cheese, freshly cracked black pepper, and pasta into a remarkably sophisticated meal despite using only a handful of ingredients.

Another local favorite is Carbonara, traditionally prepared with eggs, guanciale, pecorino cheese, and black pepper. Authentic Roman carbonara contains neither cream nor garlic—a detail proudly defended by local chefs.

Professionals seeking memorable business dinners should also sample Amatriciana, featuring tomato sauce, guanciale, and pecorino, or Saltimbocca alla Romana, tender veal prepared with sage and cured ham. Those interested in experiencing traditional Roman comfort food often order Suppli, crisp rice croquettes filled with mozzarella, making an excellent light lunch between meetings.

Coffee culture deserves equal attention.





Unlike many countries where coffee accompanies extended work sessions, Italians typically enjoy espresso while standing at the bar, often finishing within minutes. This brief ritual provides an opportunity to exchange greetings, discuss business informally, or simply pause before continuing the day.

Similarly, the Italian aperitivo tradition encourages colleagues to meet before dinner over light refreshments, offering a relaxed environment for networking and conversation. These seemingly small customs reveal much about Italian business culture: relationships matter, hospitality matters, and quality always outweighs speed.

Where Business Travelers Stay

Accommodation in Rome reflects the city's diversity, offering everything from internationally recognized luxury hotels to boutique properties housed within centuries-old buildings.

Executives attending conferences often choose hotels near the EUR district, Rome's modern business quarter, or close to the city's principal transportation hubs.

Those seeking a more immersive cultural experience frequently stay in the historic center, where many luxury hotels occupy beautifully restored palazzos featuring frescoed ceilings, marble staircases, and elegant courtyards.

Business amenities have evolved considerably in recent years. High-speed internet, executive lounges, co-working spaces, private meeting rooms, and concierge services are now standard offerings among premium hotels.

Many properties also feature rooftop terraces overlooking the city's skyline, providing exceptional venues for informal business discussions or evening networking events.

For longer stays, serviced apartments have become increasingly popular among consultants, project managers, and executives overseeing extended assignments. These accommodations offer greater flexibility while allowing visitors to experience everyday Roman life beyond traditional tourism.

Exploring Rome's Distinctive Neighborhood

One of Rome's greatest pleasures lies in discovering the unique personality of its neighborhoods. Each district tells its own story and offers a different perspective on the city.

Centro Storico: The Historic Heart

Rome's historic center remains the city's cultural and architectural centerpiece. Narrow cobblestone streets weave between churches, fountains, boutiques, cafés, and centuries-old buildings, creating an environment where nearly every corner reveals another remarkable discovery. Walking remains the best way to experience this area.

Distances between major landmarks are surprisingly manageable, allowing business travelers to transform short breaks into memorable explorations.

Whether heading toward the Pantheon, stopping for espresso in a quiet piazza, or simply observing local life unfold from an outdoor café, the historic center rewards curiosity rather than rigid itineraries.

Trastevere: Rome's Bohemian Soul

Across the Tiber River lies Trastevere, one of the city's most beloved neighborhoods.

Its ivy-covered buildings, narrow alleyways, artisan workshops, and lively squares offer a distinctly different atmosphere from the grand avenues surrounding Rome's principal monuments. During the day, Trastevere feels relaxed and residential. Independent bookstores, family-run bakeries, local markets, and small galleries encourage leisurely exploration.

As evening arrives, the neighborhood transforms into one of Rome's premier dining destinations. Restaurants spill into the streets, musicians perform in public squares, and conversations continue late into the night. It is an ideal setting for informal business dinners or post-conference gatherings that emphasize relationship-building over formality.

Monti: Creativity Meets Tradition

Located between the Colosseum and Via Nazionale, Monti has emerged as one of Rome's most creative districts

Independent designers, vintage boutiques, artisan workshops, specialty coffee shops, and contemporary galleries coexist alongside traditional businesses that have served local residents for generations.

Monti attracts entrepreneurs, designers, academics, and young professionals seeking a more contemporary side of Rome while remaining immersed in its historical surroundings. The neighborhood offers excellent opportunities to purchase locally made products, from handcrafted leather goods and ceramics to jewelry and home décor created by independent Italian artisans.



Shopping: Celebrating Italian Craftsmanship

Italy's reputation for craftsmanship extends far beyond fashion, although Rome certainly excels in that regard.

The city's premier shopping avenue, Via dei Condotti, showcases internationally renowned luxury fashion houses alongside distinguished Italian brands. Nearby streets feature elegant boutiques specializing in footwear, leather accessories, tailoring, and jewelry.

Yet some of Rome's most rewarding shopping experiences occur away from luxury storefronts.

Family-owned workshops continue to produce handmade gloves, custom stationery, silk ties, ceramics, fountain pens, and leather journals using techniques refined over generations. Purchasing these items supports local artisans while providing meaningful souvenirs that reflect Italy's enduring commitment to quality.

Business travelers often appreciate these handcrafted gifts when selecting presents for colleagues, clients, or family members upon returning home.

Museums Beyond the Vatican

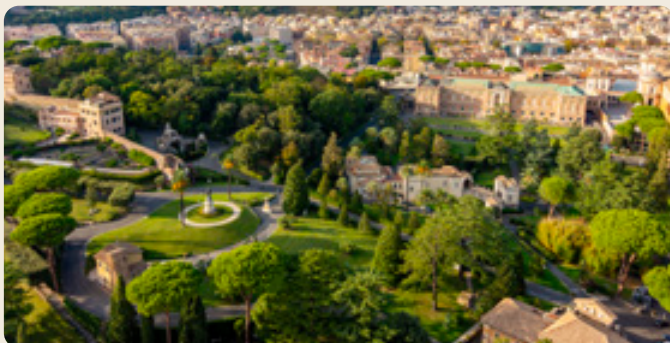
Although the Vatican Museums attract worldwide attention, Rome's cultural offerings extend far beyond Michelangelo's masterpieces.

Art enthusiasts should consider visiting the Galleria Borghese, whose collection includes extraordinary sculptures by Bernini and paintings by Caravaggio, Raphael, and Titian. The museum's relatively intimate scale allows visitors to appreciate each work without the overwhelming crowds found elsewhere.

Equally compelling is the Capitoline Museums, widely regarded as the world's oldest public museum. Located atop Capitoline Hill, its collections explore Rome's evolution from ancient civilization to modern capital through sculpture, archaeology, paintings, and historical artifacts.

Those interested in contemporary art may enjoy the MAXXI – National Museum of 21st Century Arts, whose striking architecture reflects Rome's ongoing engagement with innovation and creative expression.

Together, these institutions demonstrate that Rome's cultural identity continues to evolve rather than existing solely within its ancient past.



Gardens and Green Spaces

Business travel often involves long meetings, crowded conference venues, and demanding schedules. Rome offers numerous green spaces where visitors can recharge between professional commitments.

The expansive Villa Borghese provides a welcome contrast to the city's bustling streets. Its landscaped gardens, walking paths, fountains, and panoramic viewpoints create an atmosphere conducive to reflection or informal conversation.

Nearby, the picturesque Pincio Terrace offers one of the city's finest panoramic views, particularly during sunset when the rooftops glow beneath warm Mediterranean light.

Many professionals choose this setting for an evening stroll after conferences, finding that the peaceful surroundings provide an ideal transition between work and leisure.

Evening Rome

As daylight fades, Rome reveals another dimension of its character.

Illuminated monuments transform familiar landmarks into dramatic nighttime spectacles. The Colosseum glows against the evening sky, fountains shimmer beneath carefully positioned lighting, and piazzas fill with residents enjoying the cooler Mediterranean air.

One of the city's most enchanting experiences involves simply walking without a destination.

The route might begin at the Trevi Fountain, continue through quiet side streets lined with elegant boutiques, pass the Pantheon illuminated after sunset, and conclude in Piazza Navona where musicians and artists create an atmosphere unlike any other European capital.

Unlike nightlife centered exclusively around entertainment venues, Rome's evenings celebrate public spaces. Families stroll together, friends gather over gelato, and visitors naturally become participants in the city's social rhythm.

This sense of openness and community contributes significantly to Rome's enduring appeal. Rather than feeling like outsiders observing local life, visitors often find themselves welcomed into it.



Beyond Rome: Day Trips Worth Taking

While Rome offers enough attractions to fill weeks of exploration, its surroundings provide equally remarkable experiences. For professionals with an additional day before departure or a free weekend between meetings, several destinations can be reached easily from the capital.

Tivoli: Villas, Gardens, and Renaissance Elegance

Approximately 30 kilometers east of Rome lies Tivoli, a historic town famous for two extraordinary villas that showcase different periods of Italian grandeur.

Villa d'Este is celebrated for its Renaissance gardens, elaborate fountains, and carefully designed landscapes. Built in the 16th century, the villa represents the Renaissance belief that architecture and nature could be combined to create harmony and beauty.

Equally impressive is Hadrian's Villa, the vast imperial retreat constructed by Emperor Hadrian during the second century. The complex demonstrates the sophistication of Roman engineering, incorporating libraries, theaters, temples, gardens, and residential areas.

For business travelers, Tivoli offers a peaceful contrast to Rome's energy and provides an opportunity to experience Italy's history in a more relaxed setting.



Ostia Antica: Ancient Rome Without the Crowds

Often described as Rome's forgotten archaeological treasure, Ostia Antica was once the ancient port city that connected Rome with the Mediterranean world.

Today, visitors can walk through remarkably preserved streets, marketplaces, residential buildings, temples, and public baths. Unlike the Colosseum or Roman Forum, which attract millions of visitors annually, Ostia Antica offers a quieter and more immersive experience.

For professionals interested in trade, logistics, and economic history, the site provides fascinating insight into how the Roman Empire managed commerce, transportation, and international exchange.

The remains of warehouses, commercial buildings, and harbor infrastructure demonstrate that globalization and interconnected economies are not exclusively modern concepts.

Castel Gandolfo: A Peaceful Escape

Located in the Alban Hills overlooking Lake Albano, Castel Gandolfo offers one of the most scenic escapes from Rome.

Known primarily as the summer residence of the Pope, the town combines historic architecture, lake views, and excellent local cuisine. It is an ideal destination for travelers seeking a slower pace after several intensive days of meetings.

A leisurely lunch overlooking the lake, followed by a walk through the town's charming streets, provides a refreshing change from the urban environment.

Understanding Italian Business Etiquette

Successful business interactions in Rome often depend as much on cultural awareness as professional expertise.

Italian business culture places importance on personal relationships. Introductions, trust-building, and informal conversation often precede detailed negotiations. A few principles can help international professionals navigate meetings effectively:

Relationships Come First

Italians often prefer to know the person behind the business proposal. Taking time for introductions, conversation, and shared meals is not considered a distraction from business, it is part of the process.

Dress Professionally

Appearance remains important in Italian professional culture. Business attire is typically elegant and polished, particularly in industries such as finance, consulting, luxury goods, and law.

Respect Meal Traditions

When invited to a business meal, rushing through courses may be perceived as unusual. Meals are opportunities to strengthen relationships, exchange ideas, and demonstrate appreciation for Italian hospitality.

Sustainability and Responsible Travel in Rome

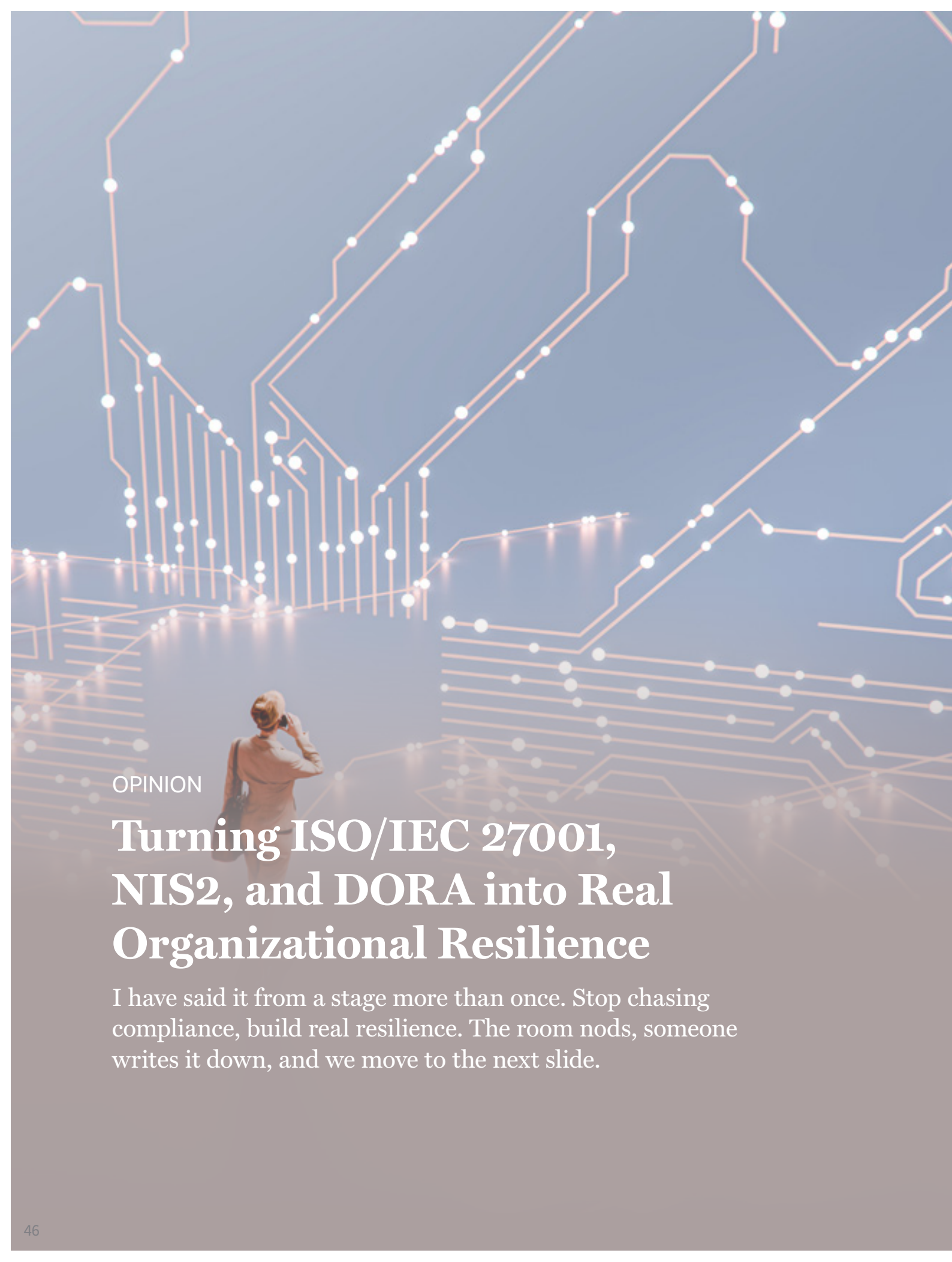
Like many major European destinations, Rome faces the challenge of balancing tourism growth with preservation of its cultural heritage.

Responsible travelers can contribute by making thoughtful choices:

- ▶ Supporting local restaurants, artisans, and independent businesses
- ▶ Choosing accommodations with sustainable practices
- ▶ Using public transportation where possible
- ▶ Respecting historical sites and local communities
- ▶ Avoiding overcrowded areas during peak hours when possible

For organizations sending employees internationally, sustainable travel policies increasingly influence destination choices. Rome offers opportunities to combine corporate travel with responsible tourism practices, particularly through walking-based exploration, local experiences, and support for traditional businesses.

By [PECB](#)

A person in a light-colored suit is seen from behind, looking out over a vast, glowing digital landscape. The background is a complex network of orange and white lines, resembling a circuit board or a data network, set against a blue gradient. The lines are illuminated with small white dots, creating a sense of depth and connectivity. The overall mood is futuristic and technological.

OPINION

Turning ISO/IEC 27001, NIS2, and DORA into Real Organizational Resilience

I have said it from a stage more than once. Stop chasing compliance, build real resilience. The room nods, someone writes it down, and we move to the next slide.

It sounds serious, the kind of line a person who has seen a few things is supposed to say. And in eighteen years I have not once watched that sentence change what an organization did the following Monday.

So what do we actually mean by it? I have started to think we do not know. We say it at every conference and put it in every board paper. It has become the polite way to look down on compliance without having to say what we would do instead. And the two words carrying all of that weight, “real” and “resilience,” are the two nobody stops to define.

Start with “real.” We only put that word in front of something when we already suspect the rest of it is not. Nobody says real gravity, or real arithmetic. When we say real resilience, we are admitting, without admitting it, that most of what we file under resilience is paper. The policy nobody reads. The recovery plan nobody has run.

The four-hour recovery target sitting next to a system that has never once been switched off to see whether the four hours hold. “Real” is us confessing, in a single word, that we do not quite trust our own vocabulary. It is the most honest thing in the sentence, and the part we skip past fastest.

Then there is “resilience,” and this is where I part ways with how most of us use it. Resilience is not vague. It is not a feeling, not a maturity score, not a color on a heat map. It is what is still working when something breaks, and how long it takes to come back. That can be measured.

Recovery time is a number. The share of a service still running while you are under attack is a number. And the figure almost nobody writes down, the gap between the moment you know something is wrong and the moment someone actually decides what to do about it, is a number too. In most of the incidents I have seen, that is where the hours (and days) disappear. Not in the technology. In the wait for someone with the authority to say yes.

I have watched a company lose a day and a half to exactly that. On paper everything was in order. ISO/IEC 27001 certified, NIS2 measures approved by the board, a continuity plan with a recovery target written next to every critical system. Then something actually broke. The failover took twenty minutes. Getting a human being with the authority to approve it took eleven hours, because the one person who could say yes was travelling and nobody below them would own the call. The plan worked perfectly. The recovery took thirty-six hours. And for all thirty-six of them, on paper, they were resilient.

That is the gap the phrase keeps pointing at without wanting to stand inside it. So the honest question is not whether these frameworks make you resilient. It is where, precisely, each of them touches that gap, because they do not touch it evenly, and most of what they ask for does not touch it at all.

ISO/IEC 27001 gets treated as the paperwork standard, and done lazily it is one. But there is a moment in it that is real, and it is not in Annex A. It is in the cycle: you assess a risk, you decide how to treat it, you act, you check. The whole thing lives or dies on the word “decide.” A risk register is worth something at exactly one moment, when a person with a budget looks at a line and commits to changing it. Everything before that is preparation and everything after it is admin. When the register gets filled in, color-coded, reviewed and filed with no decision in the middle, you have not managed a risk. You have produced a signed, dated, audit-ready document proving you saw it coming and did nothing. The auditor cannot tell those two apart. The incident can.

NIS2 does something ISO never did. It drags cybersecurity out of the IT department and puts it on the board, which now has to approve the measures, oversee them, sit through training, and can be held personally liable when it all goes wrong. It also starts a clock: twenty-four hours for an early warning, seventy-two for a fuller one, and a month for the final report. Both of those are real. A named, liable person decides differently from a committee, because they cannot make an awkward choice disappear by passing it around the table until it dissolves. And a twenty-four-hour clock cannot be answered with a binder.





And here is the uncomfortable part, the part that has nothing to do with technology. The reason we do not run the test is not that we lack the tools or the budget. It is that a test can be failed, and failed in front of the very people we spend the rest of the year reassuring. A register can only ever agree with you. A test can contradict you, on the record, with your name on the report. So we quietly choose the version that cannot embarrass us, and we call it prudence. That is a human problem long before it is a security one, and it is the real reason the frameworks stop where they do. Every framework ever written can require a plan. Not one of them can make you willing to find out whether it holds.

The first is to test the thing you are most afraid to test. Every organization has one system it will happily run a tabletop around and never actually touch, because everyone in the room quietly knows what would happen if they did. That system is your real risk, and the reluctance is the tell. DORA forces banks to do this on live production. Nobody forces a hospital, a water utility, or a logistics operator, and they are the ones with the most to learn. You do not need a regulator to schedule your own failure. You need to decide you would rather fail a test you called than an incident you did not.

But I want to be honest about what the clock does and does not do. It does not make your response any good. It shortens the time you have before you must admit you are in trouble. It does nothing to get you out of trouble faster. You can file a flawless warning at hour twenty-three and still be down for nine days. NIS2 measures how quickly you own up. It says nothing about how quickly you recover.

DORA is the only one of the three that refuses to take your word for it. For the entities it covers, a plan is not enough. You test the systems behind your critical functions every year, and if the regulator considers you significant, you sit through threat-led penetration testing, on your live production systems, run by people paid to get in, with your own defenders kept in the dark until it is over.

That is the closest anything in European regulation comes to measuring resilience honestly, because it does the one thing a document cannot: it breaks something on purpose and watches what you actually do. A report from that kind of exercise does not tell you that you have a plan. It tells you what the plan did when someone competent went looking for the difference between your architecture diagram and your Tuesday. There is always a difference. DORA just makes you find it before an attacker does.

Line the three up and the same thing shows through all of them, even though only one says it out loud. Each produces resilience where it forces a decision, an owner, or a proof, and produces paper everywhere else. Which means the work that matters is not implementing the framework. It is the handful of things the framework will never make you do, done anyway.

A real test does the one thing a register cannot. It finds the failure you never wrote down: the runbook that quietly assumed an engineer who left in the spring, the critical function that turned out to rest on a single supplier nobody in the room owned, the recovery target that held on paper and became half a day the moment someone had to be woken to approve the fix.

None of that lives in a risk register, because a register holds what you already know, and an incident is built almost entirely out of what you assumed. A tabletop where everyone reads the plan aloud will not surface it either; people are far too reasonable in a meeting room. You find it by pulling the plug, in a window you chose, and watching what the organization does when the thing it was promised could never fail, fails. That is the difference between believing you can recover and knowing it.

The second is to measure the wait, not only the recovery. Put a stopwatch on the gap between detection and decision, in every drill and every real incident, and write the number down. It is the most useful figure you are not collecting, and it is almost never a technology number. It is a governance number. It tells you whether anyone is actually allowed to act at three in the morning, or whether your entire response is quietly waiting on one person's phone. Track it across a year and it stops being an anecdote and becomes a trend, and the trend is nearly always the same: the technical half of your response has been optimized to death, and the decision sitting in front of it has never once been touched. That is the number I would take to a board, because it is the one thing they can actually fix, and usually the one thing they are the cause of.

Which leads to the third, the one that costs nothing and gets skipped anyway: decide who can say yes before the night it matters. Name the person. Name the deputy. Agree in daylight, and in writing, the irreversible calls they are allowed to make on their own, cutting a system off, failing over, going public, so that nobody burns eleven hours hunting for permission while the clock NIS2 handed you runs down. This is what NIS2's liability is actually good for. Do not use it only to punish the board after the fact. Use it to make the board delegate the decision before the fact.

The fourth is smaller and blunter. Restore a backup you actually depend on, on an ordinary day, with no warning. A backup that has never been restored is not a backup. It is a hope with a schedule.

None of that earns a certificate, and none of it is a maturity model. All of it produces the one thing no framework can hand you: evidence that the plan works when nobody is watching it. It is slower and less impressive than a certificate, and it is the only version that survives contact with a real incident.

So here is what I would put on the slide, instead of “real resilience, not compliance.” A question. When did you last break your own critical service on purpose, who was in the room to watch it fail or come back, and how long did it take before someone was allowed to decide? Answer that and you already have the thing; you never needed the slogan. Fail to answer it and no framework is going to hand it to you, because the frameworks were only ever an invitation to decide, to own it, and to prove it.

Proving it is the part almost everyone skips, quietly, while certifying the standard, appointing the owner and filing the report inside the hour. You do not find out whether you were resilient on the day the certificate arrives. You find out at three in the morning, in the time it takes one person to pick up the phone and say run it.



Christophe Mazzola

GRC Lead at Cresco Cybersecurity and
Founder at Cyber Academy

Christophe Mazzola is a security visionary and GRC Lead at Cresco Cybersecurity, on a mission to simplify complex IT. Creator of the Cyber Academy and author of the book “Être en Cybersécurité: a cybersecurity guide for everyone,” he’s served diverse sectors—from SaaS to aerospace—and notably led 100+ audits for Wallonia’s public administration. He’s a PECB, ISACA, and TEDx speaker with a knack for translating cybersecurity into plain language, reflecting his motto: “If you can’t explain it simply, you don’t understand it well enough.” For more, visit: www.celasconsulting.eu.



Why Every Executive Should Understand Technology Risk?

Technology risk is not an issue that c-suite executives can delegate entirely to their IT Department.

In our digital business environment, any type of technology risk such as a technology software or hardware failure, cyber threats, AI-related risks, data privacy breaches, third-party exposure, and digital disruption can directly affect a company's revenue (aka bottom line), reputation, regulatory compliance, customer trust, and organizational survival. Therefore, every executive needs enough technology-risk literacy to ask the right questions, challenge assumptions, make informed decisions, and ensure that the company's technology actually supports, and not threatens, the organization's strategic objectives.

Technology risk is now multifaceted and specific for each company based on their products, services and Information Technology (IT) infrastructure. Nonetheless, technology risk is now part of the overall wider business risk. It must be part of the overall Enterprise Risk Management (ERM) conversation.

Executives make decisions daily about different technology risk such as Digital Transformation, cloud adoption, Artificial Intelligence (AI), cybersecurity investments, customer-facing platforms, data, outsourcing and third-party providers, business continuity, and new technology-enabled products and services which the company has to interact with. These are business decisions that come with technology-risk consequences.

The cost or financial impact of any type of technological failure will go far beyond the IT Department and impact the entire business, such as reputational damage, online and in-person sales, potential breach of laws and regulations, etc. Now, an executive does not necessarily need to become a technologist (unless that's the company's business model), but they must understand how technology can create or amplify business risk.

What Every Executive Should Actually Try to Understand?

Although the executive of any organization, may not be an expert in technology or need to become cybersecurity professional, there is an extremely important need to have technology-risk literacy.

As a guide, here are some questions which they can ask themselves and pose to their IT team.

1. What technology risks could prevent us from achieving our strategic objectives?
2. What are our most critical systems and data?
3. What happens if one of those systems becomes unavailable today?
4. How exposed are we to cyber-attacks and technology failures?
5. What risks are created by our third-party technology providers?
6. How are we governing our use of AI and emerging technologies?
7. Can we detect, respond to, and recover from a major technology incident?

The Board and Executive Team Must Govern Their Technology Risk

The discussions regarding the governance and management of technology risks should occur frequently in boardroom and management discussions, for example on a monthly basis and not only after something goes terribly wrong.

The management of technology risk starts with the leadership team and not just left to their IT department only to figure everything out.

Executives should ensure there is a clear responsibility and accountability to the various IT roles/positions, assets, and technologies across the organization. Executives must understand the financial investment on acquiring the various technologies, as well as the impact of not continuing to invest in the latest technologies.

C-suite individuals must develop an agreed upon and approved risk appetite to understand, analyze, and pursue new technologies. Be able to ensure there is regular reporting on their testing, usage, and implementation into daily operations of the company. Also, the adoption of any technology adheres to both national and regulatory compliance. There is always continuous improvement process to monitor and evaluate ways to enhance their employees' performance and the organization.

Technology Risk must be part of your overall Enterprise Risk Management (ERM)

Firstly, Technology Risk and Cybersecurity should be presented as an Enterprise Risk Management (ERM) issue. It should link to your company's overall strategy how specific technologies e.g. Artificial Intelligence (AI), Digital Analytics, Document Processing Automation tie-in to achieve the company's objectives and also what issues can arise that will cause the business to fail.



Technology Risk Can Threaten Strategic Objectives

C-suite executives are responsible for the overall strategy of their business. For example, if an organization wants to expand internationally to launch a product or service, introduce AI, migrate to the cloud or transform customer experience, technology risks need to be considered alongside the upside of the business opportunity.

So when an opportunity arises, an executive must look at the opportunity, the technology and its risk, the controls that are or can be put in place, the business continuity system and the resulting business value.

A simplistic model for executives to follow could be:

Opportunity > Technology > Risk > Business Continuity > Business Value for their customers.

Great executives will not necessarily avoid technology risk, instead they try to understand and analyze it well enough to take intelligent and calculated risks, which they can support based on the current information.

Consider Sub-Dividing Various Areas of Technology Risk

Artificial Intelligence as a Technology Risk

Organizations are rapidly adopting AI for various business processes or areas of the business such as decision-making, customer service using AI agents, marketing, and human resources for recruitment, summarizing of legal documents, and even product development.



How can we govern the technology risks posed by AI?

The question for executives is no longer simply “How can we use AI?” But, they must also understand, “How do we govern the risks created by using AI?”

- ▶ **AI Hallucinations:** Many of the AI risks include AI hallucinations. An AI hallucination is when an artificial intelligence system makes up false information and presents it as convincing, but it is entirely wrong.
- ▶ **Bias and Discrimination:** There may be a bias and discrimination risk in terms of both the test and training data in the Artificial Intelligence (AI) system. Training data is used to teach the AI model to recognize patterns and adjust its settings. Testing data is a separate dataset used to evaluate the accuracy of the information and how well the trained model performs on new information.
- ▶ **Lack of Transparency:** The real technology risk here for executives is hidden bias, when there is no accountability to blame anyone when the opaque system makes a costly mistake and erodes trust, as users do not feel safe and can verify how it works.
- ▶ **Shadow AI:** Is the unauthorized use of artificial intelligence tools or applications by employees without the approval or oversight of an organization’s IT and security teams. Workers often use these unchecked or not thoroughly examined public applications to speed up tasks, which in turn risk data leaks and compliance breaches.

Data Privacy as a Technology Risk

The c-suite of many public and private organizations do not fully appreciate how they collect, process, store, retrieve and access, dispose of data and information they use daily to make decisions and move the business forward. This data or information (whether it be digital or paper-based information) can now become a technology risk in and of itself.

The technology risk here is if employees try to destroy or sabotage the data, there is some type of unauthorized leak of information or even a cyber-attack, this can possibly cripple an organization due to a heavy reliance on its stored data and information. Many scenarios can simultaneously occur and have a domino effect.

One of the solutions executives can look at include implementing a structured approach to privacy management, such as ISO/IEC 27701 Privacy Information Management System (PIMS), rather than leaving privacy practices scattered across departments.

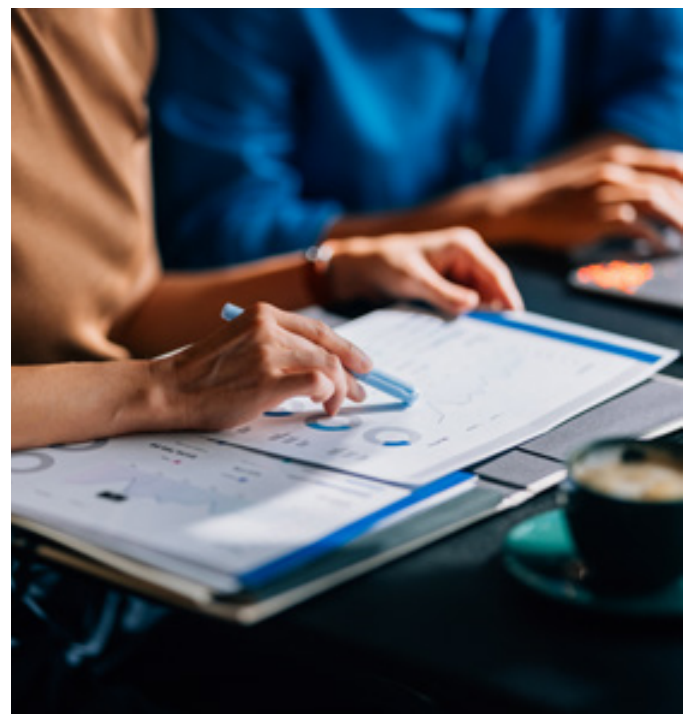
Digital Transformation and Technology Risk

As organizations continue to digitize their business processes, an increased digital dependency can create a concentrated risk and systemic vulnerabilities, if the Digital Transformation process occurs without proper risk management.



Distinction between digital transformation and digital resilience

Although Digital Transformation can make us better using the right technology, becoming a Digital Resilient company will allow the executives to ask “What happens when the technology fails?” They should be able to address that through business continuity, incident response teams, and change management processes and procedures being in place.



Technology Risk Is Also about Opportunities and (Technological) Innovation

To be clear, technology risk is not something executives should fear. Rather, leaders must understand that taking on the right risks actually enables innovation. The objective here for leaders is not to eliminate every single risk. Instead, the executives who understand, manage and consciously accept appropriate risks while pursuing opportunities are the ones who tend to make better decisions about utilizing new technologies, such as AI adoption, cloud migration, digital products, process automation, and digital transformation.

Conclusion

Organizations need to be innovative by utilizing the right technology to enhance their growth, without becoming unnecessarily vulnerable to the associated risks. The advent of new technologies will continue to transform how organizations operate, compete, and serve their customers. However, the organizations that thrive will not simply be those that adopt technology the fastest. Rather, they will be those organizations whose leaders understand the risks that accompany that specific technology and build governance, resilience and accountability measures or the right frameworks that are necessary to manage them.

Executives do not need to know how to write the code, configure the network or investigate the cyberattack. But, they do need to understand what is at stake, what questions to ask and what decisions must be made. In the digital economy, technology risk literacy is no longer a technical skill, but instead it is a leadership requirement.



Nicholas Ramsey

CEO/Consultant/ISO Trainer/IT Auditor at
N Ramsey Consultancy Ltd

Nicholas serves as the Chief Executive Officer (CEO) for N Ramsey Consultancy Ltd, which is a dynamic training and consultancy professional services company.

He is a Certified Trainer and professional in the following specialized areas: ISO 31000 Enterprise Risk Management (ERM), ISO 45001 Occupational Health & Safety Management System, ISO/IEC 27001 Information Security Management System, ISO 22301 Business Continuity Management System, ISO 9001 Quality Management System and ISO 42001 (Artificial Intelligence Management System).

Nicholas has over 20 years' of experience working at various organizations such as KPMG, AT&T, Microsoft, Apple, University of the West Indies (UWI), UWI-ROYTEC, and Deposit Insurance Corporation (DIC) to name a few, across the industries of: Financial Services, Information Technology, Auditing and Consultancy, Telecommunications, and Tertiary Education.

His degrees include a master's degree in information technology from Georgia Southern University (GSU) and a bachelor's degree in communications from Florida Memorial University (FMU) in the USA.

Additionally, Nicholas is a Certified IT Auditor (ISO/IEC 27001 Lead Auditor credential), a Member of The Institute of Internal Auditors (IIA), an Associate Member of The Business Continuity Institute (AMBCI), and possesses ITIL V3 certification.

His background includes advising companies in Project Management, Digital Transformation, Data Privacy and PCI-DSS Compliance.

He is a faculty member of the Forensic CPA Society in Miami, FL. and a former Lecturer at the University of the West Indies at St. Augustine (UWI) and UWI ROYTEC. You can reach him at: www.nrconsultancyltd.com



Privacy as a Foundation for Innovation: Protecting Health Data in the Digital Age

There is a persistent myth in digital health circles that privacy and innovation sit on opposite ends of a see-saw, that every control added to protect a patient's data is a control subtracted from the pace of progress.

A pharmacist working at the intersection of clinical practice and digital health governance has rarely found this to be true. What they have learned, repeatedly, is the opposite of the myth. Privacy is not a brake on innovation in healthcare. It is the chassis that lets the vehicle move fast without falling apart. The pharmacies, hospitals, and platforms that treat privacy as a design principle rather than an afterthought are, almost without exception, the ones that manage to scale their digital services furthest and fastest, because they are the ones patients and clinicians are willing to use.

Trust Is the Product

Every prescription a pharmacist dispenses sits on top of a chain of trust. The patient trusts that their diagnosis, their medication history, their allergies, and their genetic predispositions will be seen only by the people who need to see them, for the purpose they were shared for. That trust is not a soft, sentimental add-on to clinical care — it is the precondition for it.

A patient who does not trust the system will withhold information. They will under-report symptoms, skip follow-up, or avoid digital services altogether. In pharmacy specifically, this shows up as patients declining to use apps that could genuinely improve adherence and safety, because they are unsure who else can see their data or how it might be used against them; by an employer, an insurer, or simply a system they never consented to in any meaningful sense.

This is the uncomfortable truth that often gets lost in digital health strategy documents: the most sophisticated AI-driven medicines optimization tool is worthless if patients will not use it because they do not trust it with their data.

A pharmacist sees this daily at the counter, a patient hesitating over a digital repeat-prescription service or declining to link their records to a new app, not because the technology is flawed but because the trust architecture around it has never been made visible to them. Privacy, in this sense, is not a compliance checkbox sitting downstream of the “real” innovation. It is upstream of adoption. Get it wrong, and the innovation never gets the data, the engagement, or the legitimacy it needs to work.

This is also a commercial reality, not just an ethical one. Digital health products live or die on engagement metrics, and engagement is downstream of trust in a way that is easy to underestimate at the design stage. A platform that is technically brilliant but perceived as careless with sensitive data will see uptake stall long before any regulator intervenes, simply because patients quietly opt out. Privacy, approached this way, is not a constraint competing with the business case for innovation. It is part of the business case.

Where the Pharmacist Sees Privacy Fail

A pharmacist's vantage point is unusually practical. They are often the last person in the chain who checks a medication record before it becomes a physical dose in a patient's hand, and they are increasingly the person navigating interoperable systems; hospital records, community pharmacy systems, GP-held data, and now wearable or app-generated data, that were never designed with a single, unified privacy architecture in mind.

When those systems fail to protect data properly, the consequences are not abstract. A misdirected discharge summary can mean an allergy flag is missing at the point of dispensing. A poorly secured interface between two systems can silently corrupt a dosing field during a data migration, and nobody notices until a pharmacist queries an implausible dose. A breach of mental health or substance-use records can cause a patient to disengage from pharmacy services entirely, quietly removing themselves from a safety net that depends on continuous, trusted contact.



An access control failure that lets the wrong member of staff view an unrelated patient's record is rarely dramatic in the moment, but it corrodes exactly the kind of institutional trust that keeps patients disclosing sensitive information in the first place. None of these are hypothetical to someone who has worked across multiple care settings; they are the everyday texture of what happens when privacy is treated as an IT department's concern rather than a clinical one.

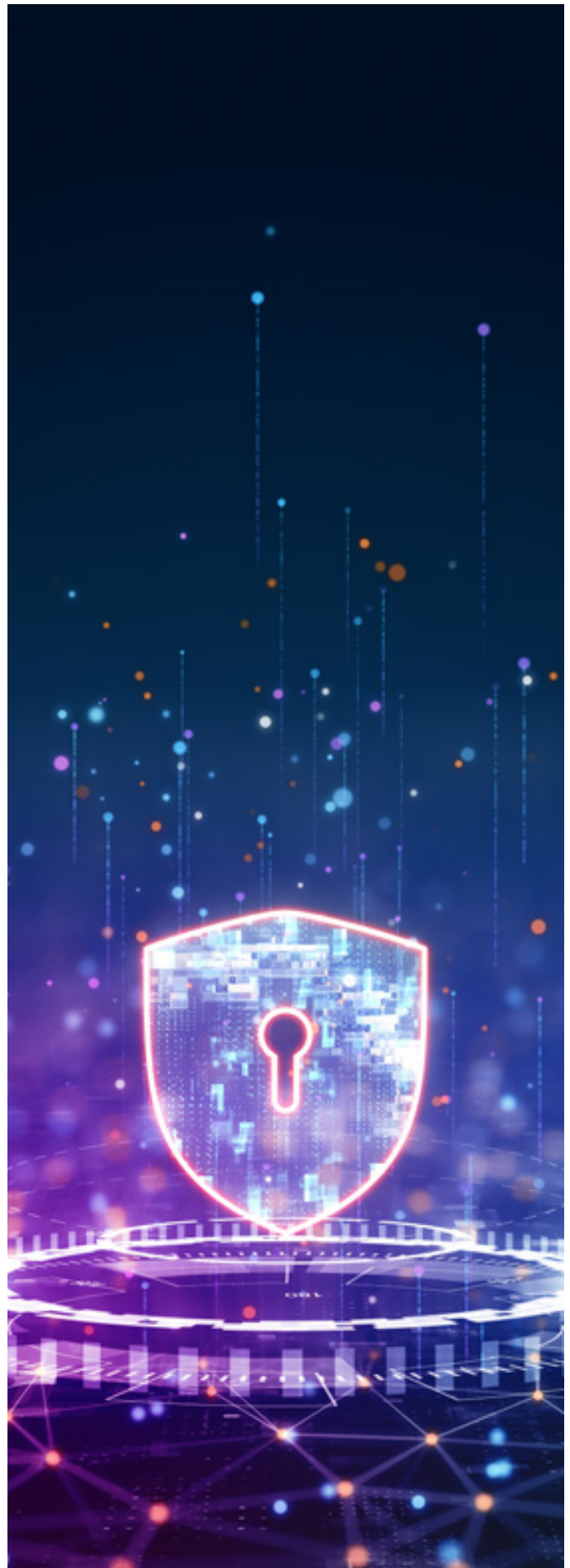
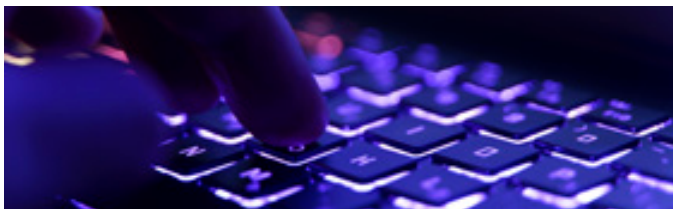
This is why, from a pharmacist's perspective, privacy failures and safety failures are very often the same failure viewed from different angles. Confidentiality, integrity, and availability of health data are not abstract information security properties, they are the conditions under which a pharmacist can trust what they are looking at before they hand over a medicine. Lose any one of those three properties and the clinical picture a pharmacist relies on becomes unreliable, whether anyone would formally classify the incident as a "privacy breach." This is also where information security standards stop being back-office paperwork and become directly relevant to frontline practice, and two standards in particular, ISO/IEC 27001 and ISO/IEC 27400, are worth examining closely.

ISO/IEC 27001: The Discipline That Makes Privacy Operational

ISO/IEC 27001 is the international standard for an Information Security Management System (ISMS) — a structured, risk-based, continually improving approach to protecting information. It is easy to mischaracterize it as a technical checklist: encryption here, access control there, an annual penetration test to satisfy an auditor. That is not what makes it valuable in a healthcare context, and it is not how a pharmacist evaluating a new digital tool should think about it either. What makes it valuable is that it forces an organization to identify its information assets, assess the realistic threats and vulnerabilities against them, and choose proportionate controls, and then to prove, on an ongoing basis, that those controls are actually working, through internal audit, management review, and continual improvement rather than a one-off certification exercise.

For a pharmacist assessing whether to recommend or adopt a new digital health tool; an app, a shared care record, a remote dispensing platform, ISO/IEC 27001 certification is a signal that turns “we take patient privacy seriously” from a marketing line into something with substance behind it. It suggests that access to patient data is governed by the principle of least privilege, that changes to systems handling patient data go through a controlled process with the ability to roll back errors, that incidents are logged and investigated, and that any third-party suppliers involved are held to equivalent standards. This matters enormously in modern pharmacy practice, where a single dispensing decision may draw on data passing through several different organizations before it ever reaches the pharmacist’s screen, each one a potential point of weakness if it is not held to the same bar as the rest.

The link back to innovation is direct. New digital pharmacy services, such as automated repeat ordering, AI-supported medicines optimization, integrated care records, depend on data flowing smoothly between systems and organizations that did not necessarily build trust with each other from scratch. ISO/IEC 27001 provides a shared, internationally recognized language for that trust. A digital health start-up that can demonstrate ISO/IEC 27001 certification does not need to convince a pharmacy’s information governance lead of its security posture from a standing start; it has already done the structural work. That shortens the path from pilot to adoption, and it reduces the chance that a rushed integration introduces an error a pharmacist later must catch, or worse, does not.

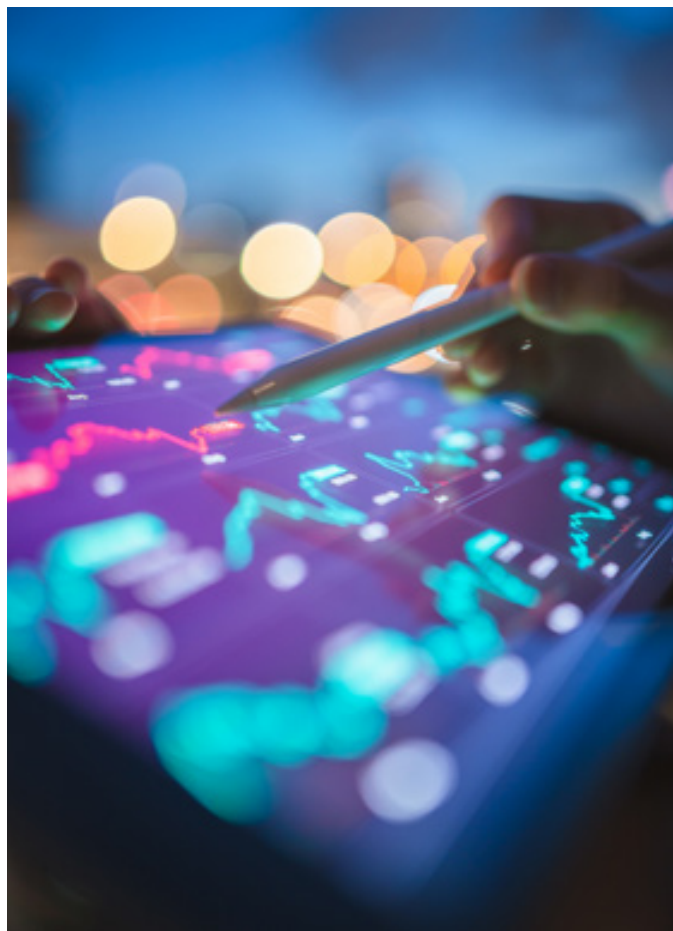


ISO/IEC 27400: Privacy in a World of Connected Things

If ISO/IEC 27001 addresses the management of information security broadly, ISO/IEC 27400 addresses something more specific, and increasingly unavoidable in modern pharmacy practice: security and privacy for Internet of Things (IoT) systems. Health data today rarely originates from a single, contained clinical record. It increasingly comes from connected devices; continuous glucose monitors, smart inhalers, connected blood pressure cuffs, medication adherence trackers, that feed data back into the systems a pharmacist relies on, often in near real time, and often before a clinician has had any opportunity to sense-check the reading against the wider clinical picture.

For a pharmacist, this matters in a very concrete way. Adherence data from a connected inhaler or a smart pill dispenser can directly inform a medicines review; a corrupted or spoofed reading from a compromised device could lead to an incorrect clinical conclusion just as easily as a paper record error could, but with far less opportunity to spot the anomaly before it influences a decision. ISO/IEC 27400 provides guidance specifically tuned to this landscape: how to manage the security and privacy risks arising from the scale, variety, and physical exposure of connected devices, many of which have limited capacity for traditional security controls and are deployed in genuinely uncontrolled environments; a patient's home, their pocket, their bathroom cabinet, rather than a hospital server room where physical and network access can be tightly governed.

The innovation link here is, again, not incidental. The entire value proposition of connected adherence and monitoring tools in pharmacy is trust at scale: thousands of devices, continuously transmitting sensitive data that a pharmacist may act on directly, often without the reassurance of a second, independent data point. ISO/IEC 27400 gives the manufacturers and platforms behind these devices a structured way to build that trust in from the outset, rather than a pharmacist discovering, after the fact, that “connected” and “secure” were never quite the same thing for a given product.



Bringing the Two Together

What stands out to a pharmacist moving between the dispensary and the wider digital health landscape is that ISO/IEC 27001 and ISO/IEC 27400 are not competing frameworks but complementary layers. ISO/IEC 27001 governs an organization's overall approach to information security — policies, risk assessments, governance structures, supplier oversight. ISO/IEC 27400 zooms in on the specific and rapidly growing category of risk introduced by connected devices, which increasingly sit at the very point where health data is first captured, before it ever reaches a managed system covered by a conventional ISMS.

Together, they describe a defensible position: an organization that can demonstrate both the management system governing information security, and the specific controls addressing the devices generating an ever-larger share of the data a pharmacist now works with. That combination is what allows a pharmacy-facing digital innovation to move from pilot to routine use without a pharmacist, or a governance board, discovering late in the process that nobody has properly mapped where the data goes, who can touch it, and what happens if a device is compromised. Treated separately, each standard leaves a gap the other was designed to close; treated together, they cover the data lifecycle from the point of capture on a patient's own device through to its use in a clinical decision.

Why This Reframes Privacy for Anyone Building in Health

The habit worth encouraging among clinicians, engineers, product leads, and boards alike is to stop treating privacy as a constraint applied to innovation after the fact, and to start treating it as one of the design requirements that makes the innovation viable in the first place. A medicines optimization algorithm trained on data patients did not meaningfully consent to share will not survive regulatory scrutiny or public trust, however accurate it is. A connected adherence platform built on devices with inadequate security will not scale past a pilot, however, elegant its clinical logic. The privacy and security architecture is not a wrapper placed around the innovation. It is part of what makes the innovation safe to put in front of a patient at all.

This is the lens a pharmacist brings, almost by professional habit, to any new digital tool that crosses their path: the question is never simply “does this work clinically?” It is “does this work clinically, safely, and in a way that respects and protects the people whose data makes it work?” ISO/IEC 27001 and ISO/IEC 27400 do not answer that question on their own; no standard does the thinking for an organization, but they provide a rigorous, internationally recognized structure within which to answer it properly. In a field where the cost of getting it wrong is measured in patient harm and eroded trust, that structure is not bureaucracy standing in the way of progress. It is the foundation progress has to stand on.

For any organization building the next generation of pharmacy and health technology, the practical takeaway is not to add a privacy review at the end of a development cycle, but to bring information security and clinical safety expertise into the room at the point a product is first conceived. Asking where the data will come from, who will be able to see it, what happens if a device is lost or compromised, and how a pharmacist will be able to trust what they are shown, are not questions to be answered retrospectively for an audit. They are the questions that determine whether the product will ever be trusted enough to matter. Standards, such as ISO/IEC 27001 and ISO/IEC 27400 do not replace this thinking; they give it a rigorous shape, a shared vocabulary across organizations, and a way of demonstrating, credibly and repeatedly, that the thinking has been done.



Paul Eversley

CEO at KPN Consultancy

Paul (MPharm, MBA Cybersecurity Management, CAIP) is CEO of KPN Consultancy, a UK-based cybersecurity, artificial intelligence, and enterprise risk management firm. He specializes in AI Management Implementation (ISO/IEC 42001), Information Security Management (ISO/IEC 27001), Enterprise Risk Management (ISO 31000), policy creation, and risk assessments across financial, insurance, healthcare, retail, and hospitality sectors.

His expertise has assisted pharmaceutical organizations with enterprise risk management services. As a GPhC registered pharmacist, he campaigns for GDPR and data protection awareness in healthcare. Paul serves as a Digital Health AI and Innovation Fellow and Clinical Safety Officer (CSO), where he advances the safe and effective integration of AI technologies in clinical settings.



THE STANDARD

How Quantum Computers Work

At first glance, a quantum computer doesn't resemble the devices we're used to. Picture a shimmering chandelier of wires, suspended inside a vacuum-sealed chamber and cooled to temperatures just above absolute zero. It hums softly, not with the whirl of fans or blinking LEDs, but with the eerie strangeness of quantum physics.

This is the world of quantum computing – not about faster chips, but a radically new way of processing information. To understand how quantum computers work, we need to move beyond raw performance and explore their fundamental principles: **how they use qubits and the laws of quantum mechanics** to solve certain problems dramatically faster than classical machines.

In this article, you'll find the quantum computer explained from the ground up. We'll unpack the basics, look at how quantum processors compute, and explore the different types of systems in development today.

The goal is clarity, not hype: a guided tour of what makes these machines so radically different from the laptop on your desk.

Quantum Computer Basics

Every computer, no matter how advanced, begins with the same building block: the bit. It's almost unbelievable that everything your laptop does – from streaming films to crunching spreadsheets – can be broken down into **strings of 0s and 1s**. And those numbers aren't abstract. In hardware, they're stored as physical states: tiny charges on capacitors, magnetic spots on a disk, or patterns etched into silicon.

Quantum computers start from the same basic principle: information must be stored in physical systems. But, here, the behavior of those systems is **governed by the laws of quantum mechanics** rather than classical physics. The basic unit is the quantum bit, or qubit. Like a classical bit, a qubit produces either 0 or 1 when it is measured. But unlike its classical cousin, it can occupy a far richer set of states, mathematically represented as a vector in a **two-dimensional space** combining the basis states $|0\rangle$ and $|1\rangle$.

You'll notice we've introduced mathematical notation here. In quantum computing, the equations aren't just a way of describing what happens – they are the thing itself. Linear algebra tells us exactly how qubits behave, evolve and interact. Without it, we'd have no way to explain what makes quantum computation so different from the digital machines we know.

Superposition: Beyond 0 and 1

At first glance, qubits may seem no different from classical bits: they have two basic states, $|0\rangle$ and $|1\rangle$, which play the same roles as 0 and 1 in ordinary computing. But here's the crucial difference: qubits aren't limited to those states. They can also exist in superpositions – combinations of $|0\rangle$ and $|1\rangle$. For instance, a qubit might be partly in $|0\rangle$ and partly in $|1\rangle$, described mathematically as something like $0.6|0\rangle + 0.8|1\rangle$. The numbers in front – here 0.6 and 0.8 – are called amplitudes: they capture the “weights” of each state, and when squared, they give the probabilities of measuring the qubit as 0 or 1.

In practice, the basis states $|0\rangle$ and $|1\rangle$ often correspond to very concrete physical configurations, such as different charge levels or a photon being in one of two locations. These are easy to picture, much like the physical ways classical bits are stored – a magnetic mark on a disk, or a hole in a punch card. Superpositions, however, are harder to imagine. They're not simply "in both states at once", although that's a common shorthand. Instead, **a qubit can exist in a continuous range of states**. This ability to carry and process combinations is one of the core reasons quantum computers can do things that classical machines cannot.

Quantum Circuits in Action

A single quantum gate can do remarkable things, but real computation arises only when they're combined. Linked together, **these gates form quantum circuits**, each one nudging qubits further along a path toward a reliable answer. On paper, this may look straightforward. In practice, though, it's one of the hardest things to achieve. Quantum states are extremely fragile, and the slightest disturbance – a flicker of heat, a stray photon – can destroy them. These random disturbances are known as quantum noise.

To preserve their state, qubits must be as isolated as possible. Some particles interact so weakly with their surroundings that they are almost impossible to disturb. Neutrinos, for instance – ghost-like particles that barely interact with matter – can pass through thick walls of lead untouched. Yet that very "aloofness" makes them useless as qubits: if we can't reliably interact with them, we can't prepare, manipulate or measure their states.

This highlights the **central dilemma of quantum hardware**: qubits must be isolated enough to preserve their state, but accessible enough that we can control them. Building a functional quantum computer is ultimately about striking that balance.

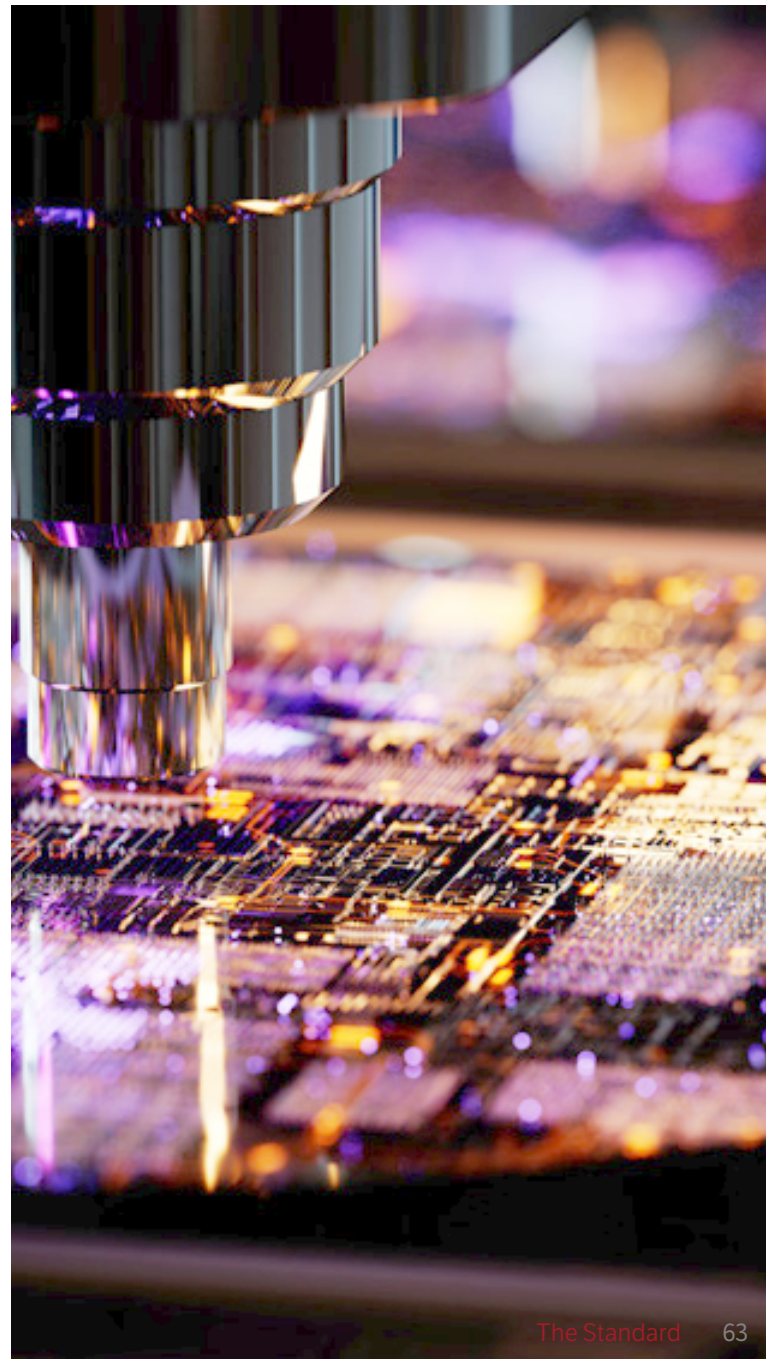
Quantum Gates: How We Program Qubits

Once you have qubits, the next question is: How do you do something with them? Superposition is powerful, but on its own it's just potential. To actually compute, we need ways to nudge, flip or blend qubits into new states. That's what quantum logic gates are for.

Think of them as the moves in the game of quantum computing. A quantum gate is a **precise operation that transforms a qubit**, or sometimes several qubits at once. They're the building blocks of every quantum algorithm, just as classical gates like AND, OR and NOT sit at the heart of everything your laptop does.

Some gates are simple. The quantum version of NOT, for instance, flips $|0\rangle$ into $|1\rangle$ and vice versa. But qubits can also exist in combinations of those states. If a qubit is in a superposition, say $|0\rangle + |1\rangle$, the NOT gate swaps the amplitudes: it becomes $|1\rangle + |0\rangle$. In other words, the weights attached to $|0\rangle$ and $|1\rangle$ simply trade places.

Other gates are more surprising. **The Hadamard gate**, for example, doesn't just flip a qubit – it blends its states into a superposition. With it, you can move from the ordinary "either/or" world of bits into the far stranger "both at once" world of qubits. That's powerful because it opens **new computational pathways**. Where a classical NOT gate is like walking back and forth along a road, the Hadamard gate is like taking a boat straight across the water. Suddenly, new routes appear, and with them, new ways of solving problems that classical machines can't match.



Making Qubits Talk

Designing circuits is one thing, but it's meaningless unless we can extract a result. That's where measurement comes in. Imagine someone hands you a qubit in some unknown state. Can you simply "peek inside" to see how it was set up? Surprisingly, no. It's not that the qubit is "being secretive", it's that quantum mechanics doesn't allow us to access the full state directly. When you measure, you only ever get one clear outcome.

To get that outcome, the qubit has to be pushed into contact with the classical world. In practice, that means letting it interact with a detector, a laser setup or an electronic circuit. The exact setup varies depending on the technology: lasers reveal the state of trapped ions, microwave circuits pick up signals from superconducting qubits, and photon detectors register which path a photon has taken.

Whatever the platform, the effect is the same: the quantum state is converted into a classical signal that we can record.

How Quantum Computer Work

We've looked at the basic building blocks of quantum computing: qubits, quantum gates, circuits and measurement. But the real magic takes place when you put them together and **watch how groups of qubits interact**. This is where the working principles of quantum computers come alive.

Entanglement: Link two qubits and they no longer behave like isolated objects. Their states become intertwined, so measuring one instantly determines the correlated outcome of the other, even if they're separated by vast distances. This strange bond isn't just a physics curiosity; it's a powerful resource that allows quantum computers to tackle certain problems far more efficiently than classical machines.

Interference: As qubits evolve through a quantum circuit, their hidden waves of probability overlap, cancel or reinforce one another. Well-designed algorithms exploit this effect to suppress wrong answers while amplifying the right ones. Grover's algorithm is a striking example, using interference to search through enormous datasets much more effectively than classical computers.

Decoherence: But there's a catch: qubits are fragile. Heat, stray particles, or even the faintest vibration, can collapse their delicate state, a problem known as decoherence. Error correction methods help, but keeping qubits stable long enough to run meaningful algorithms is one of the biggest challenges in quantum computing.

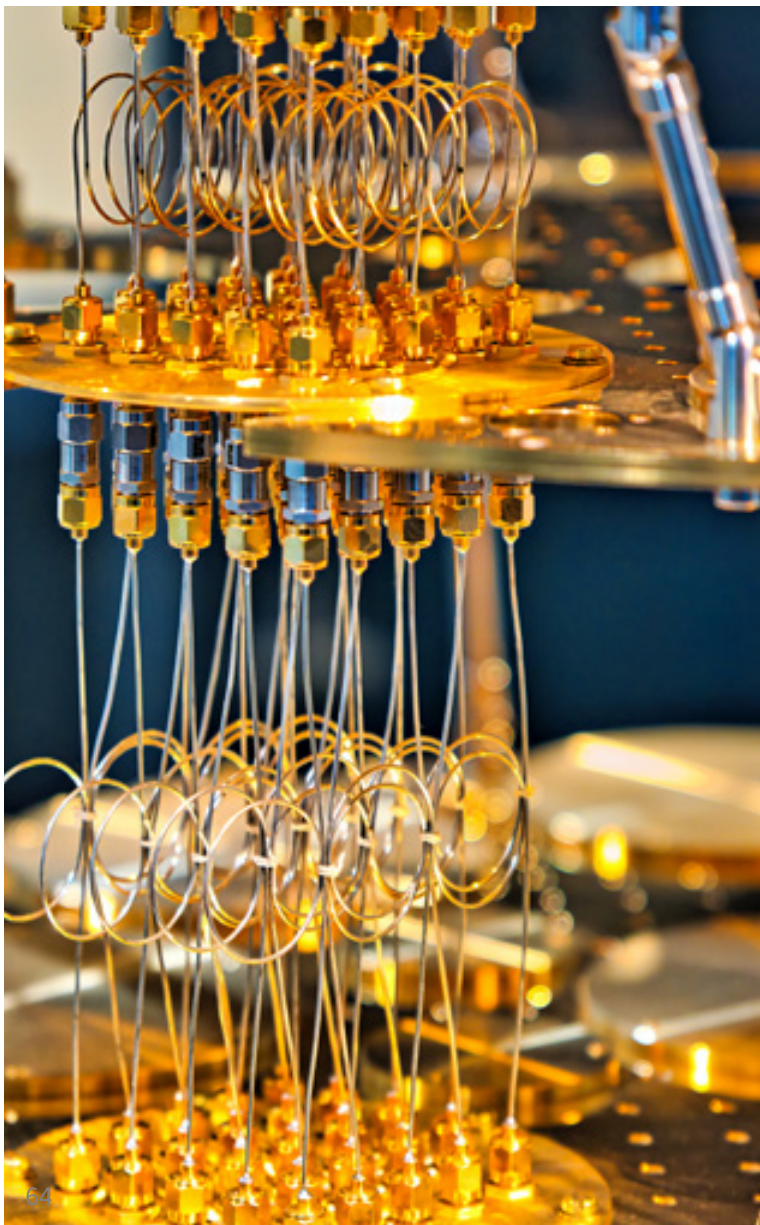
Put together, the cycle looks like this: prepare qubits in a chosen state (usually $|0\rangle$); apply quantum gates that manipulate and entangle them; design the circuit so interference amplifies the correct answers while suppressing the wrong ones; and finally measure the qubits to convert the result into classical information.

It's not just a faster way of computing, it's a fundamentally new way of thinking about what computers can do.

Different Paths to Quantum Computing

The hard part isn't the theory, it's making it work in physical machines. Around the world, researchers are testing different types of quantum computers, each based on a distinct physical system with its own strengths and limitations.

Superconducting qubit quantum computer: Used by IBM and Google, this is today's most established platform. Superconducting circuits cooled to near absolute zero switch states incredibly fast, but scaling is difficult – the more qubits you add, the harder it is to keep them stable.





Thanks to mature software like Qiskit and Cirq, developers can build, test and refine quantum programs more easily, giving these systems a natural head start.

Trapped-ion quantum computer: Companies like IonQ use lasers to trap and manipulate individual atoms. These qubits have exceptionally long coherence times, making them very stable. The trade-off is slower gate operations and greater difficulty scaling to large systems.

Photonic quantum computer: Light itself can carry quantum information, as in the models of PsiQuantum and Xanadu. Photons (light particles) resist noise naturally, but coaxing them to interact strongly enough for computation is a major challenge.

Neutral atoms quantum computer: Using optical tweezers, atoms can be arranged in grids of potentially thousands of qubits. This is a highly scalable model, though still young compared to other platforms.

Quantum dot quantum computer: This model uses nanoscale semiconductors to trap and manipulate single electrons. Built from the same materials as classical chips, they fit naturally into existing semiconductor technology. The difficulty, however, is keeping those delicate quantum states coherent at such tiny scales.

Topological quantum computer: This approach, championed by Microsoft, aims to harness anyons – exotic quasiparticles thought to form the basis of topological qubits. In theory, these qubits could be far more resistant to quantum noise than other designs, but for now the concept remains unproven.

Each of these approaches comes with trade-offs – speed versus stability, scalability versus control – and none has emerged as the definitive path forward. Together, though, they provide some of the most striking quantum computer examples, showing how diverse physical systems can be engineered to harness the underlying principles of quantum mechanics.

Standards Powering the Quantum Era

Quantum computing is no longer a distant vision, it's already here in its earliest forms. Through cloud platforms, researchers and businesses can run experiments on real machines, usually only a few dozen qubits strong. Modest as they are, these prototypes prove that the strange rules of quantum mechanics can be turned into working computation.

The challenge now is coherence, not just in qubits, but across the field itself. As designs multiply, the risk is fragmentation: competing platforms, each speaking its own language. To counter this, quantum computing has its own dedicated technical committee: [ISO/IEC JTC 3](#). Created in 2024, it brings together global experts to set common rules for how quantum systems are described, built and connected.

Standards don't hold quantum computing back – they give it direction. They provide the invisible scaffolding for interoperability and trust, turning isolated breakthroughs into a shared ecosystem. With standards, the field can move beyond dazzling experiments to become an industry with the power to transform finance, healthcare, logistics, and beyond.

► [ISO/IEC 4879:2024](#): Information technology — Quantum computing — Vocabulary

Quantum Horizon

With standards laying the groundwork, the scattered experiments of today are being shaped into a connected ecosystem. That foundation is crucial, because quantum computing is not just about faster machines, it's about reimagining what's possible. By bending the rules of physics into tools for computation, we stand at the edge of discoveries classical technology could never reach.

The road ahead will be long, full of noise, errors and technical hurdles. But history shows that every great leap in computing – from vacuum tubes to microchips to the cloud – began with fragile prototypes and bold imagination. Quantum is no different.

What makes this moment extraordinary is that we are present at the very beginning. The breakthroughs of tomorrow will be built on the questions we ask today. Quantum computing is not a distant future – it's a revolution already unfolding.

Disclaimer:

PECB has obtained permission to publish the article written by [ISO](#).

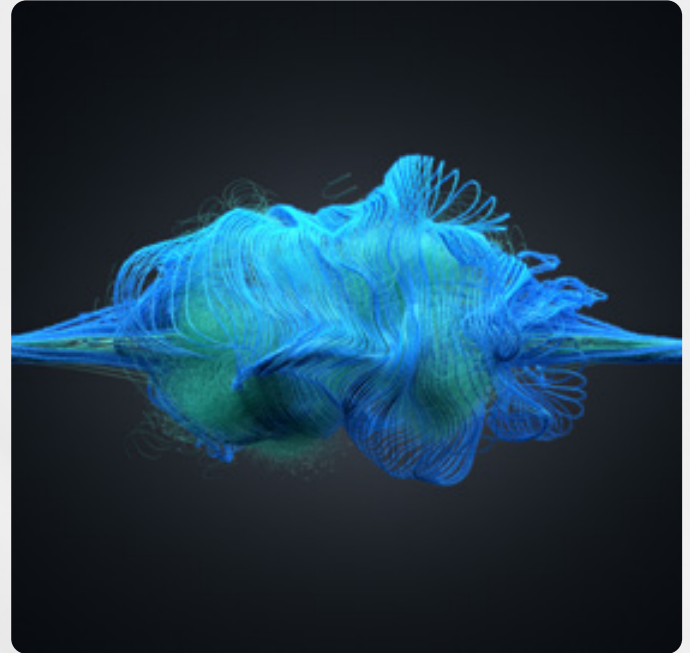
Migrating Legacy Infrastructure to

Cybersecurity teams today face a growing challenge: the sheer volume of threat data generated across networks, endpoints, cloud environments, operational technology (OT), and third-party systems. Threat intelligence feeds, SIEM alerts, vulnerability scans, phishing reports, dark web monitoring, and endpoint telemetry all produce valuable information, but without proper visibility and correlation, organizations risk becoming overwhelmed by fragmented data rather than empowered by actionable intelligence.

This is where threat intelligence dashboards become critical. A well-designed threat intelligence dashboard allows organizations to centralize, visualize, prioritize, and operationalize security intelligence in real time. Rather than forcing analysts to navigate multiple disconnected tools, dashboards provide a unified operational view of the threat landscape, helping teams identify risks faster and respond more effectively.

However, building an effective dashboard requires more than simply displaying security data on a screen. The dashboard itself must be secure, scalable, operationally useful, and aligned with the organization's security objectives.

This article outlines a practical step-by-step approach to building a secure threat intelligence dashboard suitable for modern enterprise environments.



1. Assess the Existing Environment

Every successful migration begins with visibility. Before making architectural decisions, organizations need a comprehensive inventory of their current infrastructure. This assessment should identify:

- ▶ Physical and virtual servers
- ▶ Applications and supporting services
- ▶ Databases and storage systems
- ▶ Network topology
- ▶ Identity and access management solutions
- ▶ Security tools
- ▶ Software dependencies
- ▶ Third-party integrations
- ▶ Compliance obligations

Equally important is understanding how applications interact with one another. Many legacy systems depend on undocumented connections, shared databases, or manual processes that have evolved over many years. Discovering these relationships early prevents costly surprises later in the migration.

Organizations should also classify workloads according to business criticality, sensitivity of processed data, recovery objectives, and operational complexity. Not every system requires the same migration strategy.

Build a Secure Cloud-First Architecture

2. Define Business and Security Objectives

Cloud migration should never become an IT-only initiative. Business objectives must guide technical decisions. Common goals include:

- ▶ Improving system availability
- ▶ Increasing scalability
- ▶ Reducing infrastructure costs
- ▶ Accelerating software deployment
- ▶ Supporting remote work
- ▶ Enhancing disaster recovery capabilities
- ▶ Meeting regulatory requirements

Rather than replicating existing security controls, organizations should redesign them to leverage cloud-native capabilities such as identity-centric security, automated monitoring, encryption, continuous compliance assessments, and policy-driven governance.



3. Choose the Appropriate Migration Strategy

Not every application should be migrated using the same approach. The widely adopted “6 Rs” migration framework helps determine the most suitable strategy:

- ▶ **Rehost (“Lift and Shift”)**
Applications are moved with minimal modification. This approach is fast but often fails to maximize cloud capabilities.
- ▶ **Replatform**
Applications undergo minor optimization while preserving their core architecture. This often delivers better performance without requiring complete redevelopment.
- ▶ **Refactor (Re-architect)**
Applications are redesigned specifically for cloud-native environments using containers, microservices, or serverless computing.
- ▶ **Repurchase**
Organizations replace legacy software with Software-as-a-Service (SaaS) alternatives.
- ▶ **Retire**
Applications that no longer provide business value are decommissioned.
- ▶ **Retain**
Certain systems remain on-premises because of technical, regulatory, or operational requirements.

Most organizations ultimately adopt a hybrid approach, selecting different migration strategies for different workloads.

4. Design a Secure Cloud Architecture

Security should be embedded into the architecture from the outset rather than added after deployment.

A secure cloud-first environment typically incorporates:

Identity-Centric Access Control

Identity becomes the new security perimeter. Organizations should implement:

- ▶ Multi-factor authentication (MFA)
- ▶ Role-based access control (RBAC)
- ▶ Least-privilege principles
- ▶ Single sign-on (SSO)
- ▶ Privileged Access Management (PAM)

Network Segmentation

Although cloud providers secure the underlying infrastructure, customers remain responsible for protecting their workloads. Effective segmentation includes:

- ▶ Separate production and development environments
- ▶ Private subnets
- ▶ Security groups
- ▶ Network access control lists
- ▶ Web application firewalls
- ▶ Zero Trust networking principles

Encryption Everywhere

Sensitive information should remain protected whether it is stored, processed, or transmitted. Best practices include:

- ▶ Encrypting data at rest
- ▶ Encrypting data in transit
- ▶ Secure key management
- ▶ Hardware Security Modules (HSMs) where appropriate
- ▶ Automated certificate lifecycle management



5. Migrate Incrementally

Attempting to migrate every application simultaneously introduces unnecessary risk. Instead, organizations should begin with lower-risk workloads before progressing toward mission-critical systems. A phased migration often follows a similar sequence:

1. Development and testing environments
2. Internal business applications
3. Customer-facing services
4. Core production workloads
5. Highly regulated systems

Rollback plans should exist for every migration stage, allowing organizations to restore services quickly if unexpected issues arise.

6. Automate Infrastructure Deployment

Manual configuration increases the likelihood of inconsistencies and security vulnerabilities. Infrastructure as Code (IaC) enables teams to define cloud environments using version-controlled templates rather than manual processes. Benefits include:

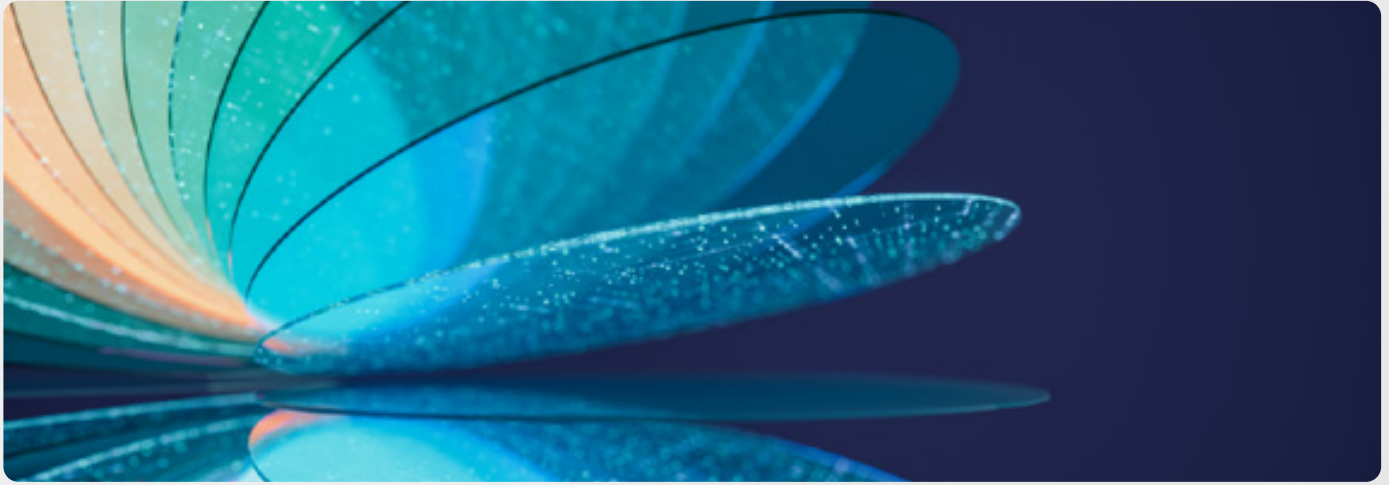
- ▶ Consistent deployments
- ▶ Faster provisioning
- ▶ Simplified disaster recovery
- ▶ Easier auditing
- ▶ Improved collaboration
- ▶ Reduced configuration drift

Automation should extend beyond infrastructure to include patch management, vulnerability scanning, compliance monitoring, and security policy enforcement.

7. Strengthen Cloud Security Operations

Migration is only the beginning. Organizations must continuously monitor their cloud environments to detect emerging threats and maintain compliance. Key operational capabilities include:

- ▶ Centralized logging
- ▶ Security Information and Event Management (SIEM)
- ▶ Extended Detection and Response (XDR)
- ▶ Cloud Security Posture Management (CSPM)
- ▶ Continuous vulnerability assessments
- ▶ Automated threat detection
- ▶ Security orchestration and response



8. Validate Performance and Resilience

Before declaring the migration complete, organizations should thoroughly test their new environment. Validation activities should include:

- ▶ Functional testing
- ▶ Performance benchmarking
- ▶ Load testing
- ▶ Disaster recovery exercises
- ▶ Penetration testing
- ▶ Backup restoration testing
- ▶ Identity and access verification

Testing should confirm not only that applications function correctly but also that security controls operate as intended under realistic conditions. Regular resilience exercises help ensure that recovery objectives can be achieved during actual incidents.

9. Optimize After Migration

Cloud adoption is an ongoing journey rather than a one-time project. Once workloads have stabilized, organizations should continuously optimize their environments by:

- ▶ Rightsizing compute resources
- ▶ Eliminating unused services
- ▶ Improving storage efficiency
- ▶ Reviewing security policies
- ▶ Updating access permissions
- ▶ Modernizing remaining legacy applications
- ▶ Monitoring cloud spending

Cloud-native architectures provide opportunities for continuous improvement that traditional data centers often cannot match.

Measuring Success

A cloud migration should ultimately deliver measurable business value. Organizations should evaluate success using key performance indicators such as:

- ▶ Infrastructure availability
- ▶ Recovery time objectives (RTO)
- ▶ Recovery point objectives (RPO)
- ▶ Security incident frequency
- ▶ Mean time to detect (MTTD)
- ▶ Mean time to respond (MTTR)
- ▶ Deployment frequency
- ▶ Infrastructure costs
- ▶ Application performance
- ▶ User satisfaction

Monitoring these metrics over time helps demonstrate return on investment while identifying opportunities for further optimization.

Conclusion

Migrating legacy infrastructure to a secure cloud-first architecture is more than a technology upgrade, it is an opportunity to redesign how IT supports the organization. By combining careful planning, phased implementation, cloud-native security controls, and continuous optimization, organizations can reduce operational risk while improving agility and resilience. The most successful migrations are those that treat security as a foundational design principle rather than a compliance checkbox. As cloud technologies continue to evolve, organizations that embrace a security-first modernization strategy will be better positioned to respond to emerging threats, meet regulatory expectations, and adapt to future business demands with confidence.

By [PECB](#)

PECB Skills

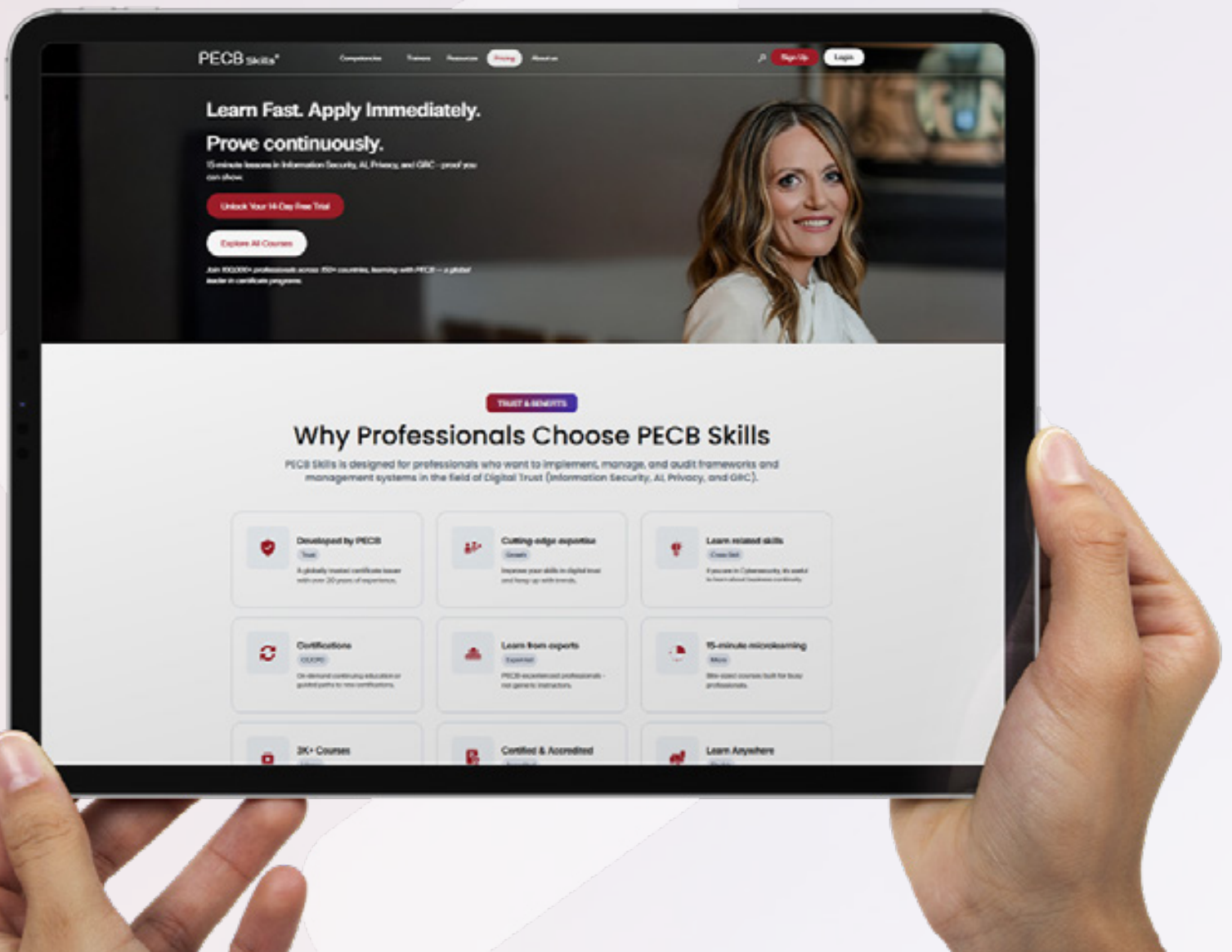
Learn Today. Lead Tomorrow.

Stay ahead in today's rapidly evolving digital landscape with PECB Skills, a flexible microlearning platform designed for busy professionals.

Access expert-led, 15-minute learning capsules covering cybersecurity, information security, artificial intelligence, privacy, governance, risk, and compliance; all available anytime, anywhere. Whether you're looking to earn CPD credits, develop new competencies, or expand your expertise through guided learning paths, PECB Skills makes continuous professional development simple and accessible.

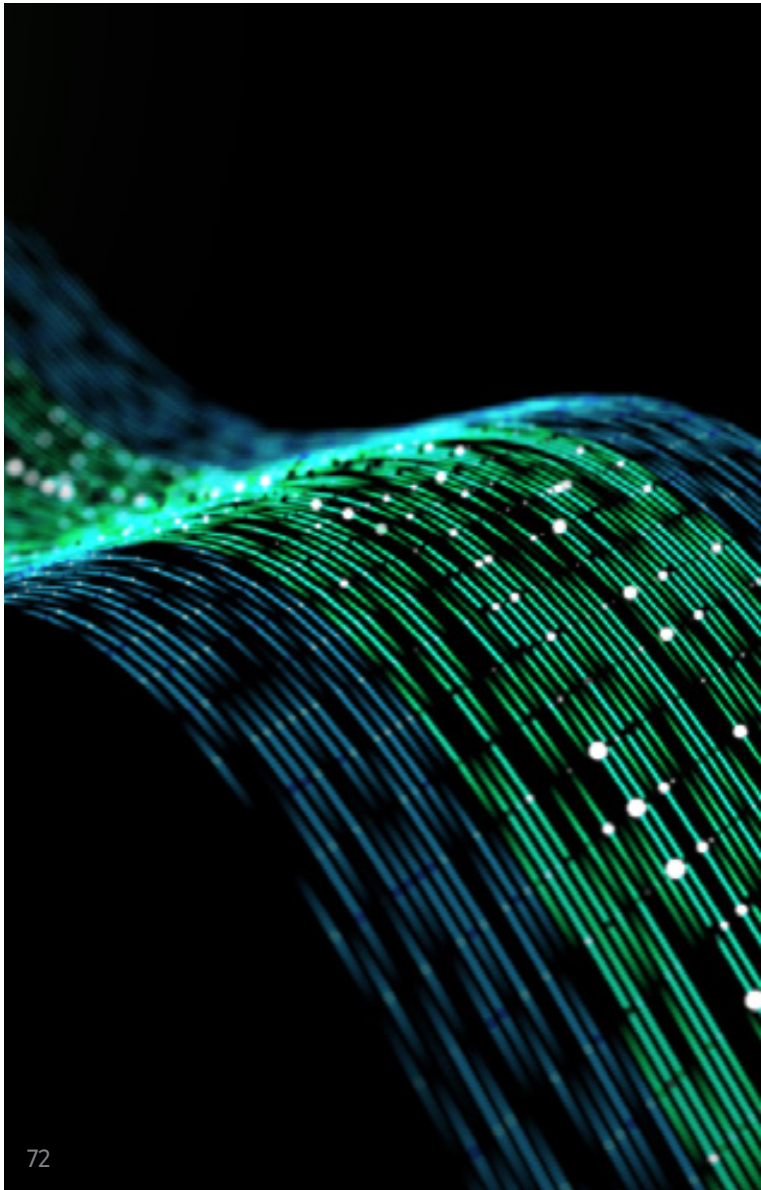
Join thousands of professionals across more than 150 countries and start building the skills that matter most. Explore PECB Skills and begin your learning journey today!

[READ MORE →](#)



Three Disciplines, One Mission: Understanding the Roles of IT, Security, and Privacy

Twenty years ago, if you asked someone in an office what “IT” meant, they would say it was the person you called when the printer jammed or your email stopped working.



Fast forward to today, and IT, security, and privacy have quietly become three of the most consequential functions inside any organization, right up there with finance and legal. And yet, a surprising number of executives still treat them as a single blob called “the tech stuff.”

That confusion is not just a semantic annoyance. It costs money, it costs trust, and in 2025 it cost the average breached organization in the United States a record 10.22 million dollars, according to [IBM's Cost of a Data Breach Report](#). Globally, the average dropped slightly to 4.44 million dollars, the first decline in five years, largely thanks to AI-assisted detection cutting the time to identify and contain a breach down to 241 days. Progress, sure. But 241 days is still eight months of an intruder having the run of the house before anyone notices.

So let's clear something up, once and for all: IT, security, and privacy are related, but they are not the same job, they do not have the same goals, and treating them interchangeably is exactly how organizations end up on the front page for the wrong reasons.

Three Disciplines, One Boardroom Headache

Here is the plain language version, the one I wish someone had handed me two decades ago.

- ▶ **Information Technology** is about making things work. Servers run, applications load, employees can log in and get their jobs done. IT's mission is uptime, efficiency, and enabling the business. If IT does its job well, nobody notices it exists.
- ▶ **Security** is about keeping the bad guys out, and just as importantly, knowing what to do when they get in anyway (because they will, eventually, no matter how good your defenses are). Security asks: who can access what, how do we detect something unusual, and how fast can we respond before “unusual” becomes “catastrophic.”
- ▶ **Privacy** is about respecting the data itself, not just protecting it from theft, but making sure it is collected, used, and shared the way the person it belongs to actually agreed to. Privacy asks a very different question than security does: even if nobody stole this data, are we using it in a way our customers, regulators, and our own conscience would be comfortable with?

Here's the part that trips people up: you can have excellent IT and lousy security. You can have airtight security and still violate privacy laws left and right, because a perfectly secured database of data you had no right to collect is still a legal landmine. These three disciplines overlap, argue, and depend on each other constantly. Treating them as one function is like asking your accountant to also be your lawyer and your therapist. They might be capable people, but that is three very different skill sets wearing one exhausted hat.

Why This Distinction Is Suddenly Urgent

For most of my career, this was an academic argument, mostly confined to conference panels and mildly heated LinkedIn threads. Not anymore. Two forces have collided to make the distinction unavoidable for anyone in a leadership seat.

The first is regulation, and it is arriving fast. As of early 2026, twenty US states now have comprehensive privacy laws on the books, with Indiana, Kentucky, and Rhode Island joining the list this year alone. States are not just passing these laws either, they are enforcing them, with an estimated 3.425 billion dollars in privacy-related fines issued across US states in 2025, [according to Gartner](#). Meanwhile, the EU AI Act reaches full enforcement on August 2, 2026, requiring organizations to maintain a documented inventory of their AI systems, a risk classification for each one, and clear disclosures about how those systems are used. If your organization cannot answer “what AI is running on our data, and why” with a straight face, you have some homework to do before that deadline lands.



The second force is AI itself, which has become both the shiny new productivity tool and a fresh attack surface, sometimes in the same product. [IBM's research](#) found that 97% of organizations that suffered an AI-related security incident had no proper access controls on their AI systems, and 63% had no AI governance policy at all. Read that again. It is not that the controls failed. In most cases, they simply were never built. That is not a security problem or a privacy problem in isolation. It is both, tangled together, and it needs both disciplines pulling in the same direction to fix.

The Practical Part; Because Nobody Wants Another Lecture

If you are a CISO, a privacy officer, or anyone with “information” somewhere in your job title, here is what this actually means for your Monday morning:

Stop letting “we are secure” and “we are compliant” be treated as synonyms in board meetings. They answer different questions, and a board that only hears one of them is flying half blind.

Build your AI governance program before regulators or auditors force you to. An inventory of every AI tool touching customer or employee data, who approved it, what data it touches, and what happens if it is wrong, is no longer optional homework. It is table stakes, and August 2026 is closer than it looks on a calendar.

Get your privacy and security teams talking to each other on a regular cadence, not just during incident response. Security tends to think in terms of confidentiality, integrity, and availability. Privacy thinks in terms of consent, purpose limitation, and individual rights. Both are right. Both are needed. Neither one alone tells the whole story of whether your organization is actually trustworthy with the data it holds.

And translate all of this into business language for your executives. Nobody in the C-suite needs to understand the difference between symmetric and asymmetric encryption. They do need to understand that a privacy misstep and a security breach are two different fires, they spread differently, and they need two different fire departments, working from the same playbook, ready before the smoke starts.



A Final Reflection

I have spent twenty years watching this field mature from “the people who fix the printer” into a set of disciplines that quite literally determine whether an organization survives a bad year. What strikes me most, looking back, is not how much the technology has changed (though it has, dramatically), but how slowly our language and our org charts have caught up to that change.

Information technology, security, and privacy are not competing priorities fighting for the same budget line. They are three different lenses on the same responsibility: earning and keeping the trust of the people whose data you hold. Get the technology right and things run. Get the security right and things stay standing when someone tries to knock them down. Get the privacy right and people actually want to keep giving you their data in the first place.

Miss any one of the three, and the other two will not save you. So the next time someone in a meeting says “let’s just get IT to handle the privacy stuff,” do the field a favor: gently correct them. Twenty years from now, I would like the next generation of practitioners to inherit a boardroom that already understands the difference, so they can spend their careers solving new problems instead of still explaining the old ones.



Carlos A. Suárez

Global Technical Success Manager,
Master of Business Administration,
MBA, CISSP

Carlos is a seasoned cybersecurity executive with over 15 years of experience, specializing in consultative sales within the cybersecurity domain. Known for a customer-centric approach and a keen focus on aligning strategies with business objectives, he is passionate about cultivating enduring relationships to maximize customer value over the long term.

His extensive background includes collaborating with global and regional enterprises in the U.S. and Latin American markets. Carlos brings a leadership-oriented mindset with attitudes and skills well-suited for management roles. He possesses a strategic vision, a commitment to fostering innovation, and a track record of successfully leading teams in dynamic, technical environments. Certified in ISO/IEC 27032 Lead Manager, ISO/IEC 27001 Lead Implementer, and ISO 31000 Risk Management, he is well-prepared to contribute to the strategic direction and technical excellence required for leadership positions in the cybersecurity landscape.



PECB Recognized with the 2026 Cybersecurity Excellence Community Choice Award

PECB has been recognized with the 2026 Cybersecurity Excellence Community Choice Award!

This recognition is especially meaningful because it is determined by votes from the global cybersecurity community, reflecting the trust and support PECB has earned from professionals, partners, clients, and its wider network.

This achievement marks an important milestone and reinforces PECB's commitment to delivering high-quality training, certification, and professional services to professionals worldwide.

Thank you to our global community for your trust and support. This achievement belongs to all of us.

[LEARN MORE →](#)

Q&A FROM WEBINAR

“The AI Governance Illusion: Why ISO/IEC 42001 and CAIP Matter More Than Ever”

This Q&A set went beyond the familiar territory of ISO/IEC 42001, addressing some of the genuinely more complex governance challenges that organizations are beginning to confront.

[Check Out the Webinar](#)



T

hese included multi-agent accountability, the limits of confidentiality in AI models, continuous auditing vs. point-in-time assurance, Shadow AI and uncontrolled usage, AI risk escalation to boards, AI governance in developing economies, human override thresholds, AI versus deterministic systems in security operations, and the future role of auditors, risk managers, and AI officers.

At its core, however, the discussion remained closely aligned with the central message of the webinar: the biggest risk is often not the absence of AI governance, but the illusion that adequate AI governance already exists.

Across the questions, several recurring themes emerged, including:

- 1. Accountability remains the biggest unresolved challenge as AI systems become increasingly autonomous.**
- 2. Human oversight is still considered essential, especially for high-impact decisions.**
- 3. ISO/IEC 42001 is an enabler, not a guarantee. Governance frameworks support trustworthiness, but they do not eliminate risk.**
- 4. Competence is becoming as important as compliance, which fits perfectly with the CAIP positioning throughout the webinar.**
- 5. AI auditing is evolving from document reviews toward behavior validation, which many participants clearly found one of the most thought-provoking topics.**

The questions collectively cover much of the practical reality organizations are currently facing when moving from ISO 27001 to ISO/IEC 42001, compliance to operational governance, human decisions to AI-assisted decisions, and static systems to agentic AI. The quality of the audience questions suggests that the webinar was attended by a highly mature and engaged audience with substantial practical governance experience.



Q: What is the biggest difference between just passing an AI audit and actually protecting sensitive data?

A: The biggest difference is that passing an audit demonstrates compliance at a specific point in time, whereas protecting sensitive data requires continuous governance throughout the entire AI life cycle. Many organizations focus on documentation, policies, and evidence collection, yet fail to address how AI systems actually behave in practice. True protection requires continuous monitoring, validation of outputs, control of training data, secure deployment practices, and ongoing oversight.

This distinction is particularly important because AI introduces new confidentiality risks that are not addressed through documentation alone. Sensitive information may be reconstructed from trained models, leaked through outputs, or become difficult to delete once embedded in model behavior. Therefore, effective protection depends not only on compliance but also on operational safeguards such as data minimization, access control, monitoring, and secure AI architectures. Passing the audit proves that governance exists on paper. Protecting sensitive data proves that governance actually works in practice.

Q: Why do many organizations focus only on getting the ISO certificate instead of actually managing AI risks?

A: The webinar describes this phenomenon as the “AI Compliance Theater.” Many organizations see certification as an objective, measurable milestone, while effective governance is a continuous and often more difficult operational activity.

As a result, management attention may shift toward demonstrating compliance rather than managing real risks.

Several factors contribute to this:

- ▶ Certification is visible to customers, regulators, and business partners.
- ▶ Governance maturity is harder to measure than obtaining a certificate.
- ▶ Existing ISO 27001 programs often create a false sense of security regarding AI-related risks.
- ▶ Organizations frequently rely on documentation instead of understanding actual model behavior and decision-making processes.

However, AI risks such as bias, hallucinations, lack of explainability, prompt injection, data poisoning, and model leakage cannot be mitigated simply by passing an audit. Governance must be embedded into operations, accountability structures, and continuous monitoring processes.

Q: The changing role of auditors raised the question of relevance as the stated new role looks more like what is equally expected of the risk officers. Kindly provide more clarity.

A: This is an excellent observation. The webinar deliberately highlights that the competencies of auditors and risk officers are beginning to overlap, but their missions remain fundamentally different. Risk officers are responsible for identifying and assessing AI risks, defining risk appetite and treatment measures, designing governance frameworks, operating ongoing risk management processes, and supporting management decision-making.

Auditors are responsible for independently assessing whether governance processes are designed and operating effectively, evaluating AI system behavior, validating evidence and controls, challenging assumptions and management assertions, and providing assurance to management and stakeholders.

Why does the role seem to converge? AI systems require understanding of model behavior, data quality, explainability, regulatory requirements, and risk management principles. Consequently, both professions need similar technical knowledge. However, risk officers manage risk. Auditors challenge whether that risk is actually being managed effectively. The auditor therefore becomes more strategic because they increasingly evaluate governance maturity, AI accountability mechanisms, and control effectiveness.

Q: How do we achieve Zero Trust AI?

A: Organizations must apply the “Never trust, always verify” principle to every component of the AI ecosystem.

Practical elements include:

1. **Verifying all inputs:** Validate data sources, control prompt inputs, and assess trustworthiness of training datasets.
2. **Verifying model behavior:** Continuously testing outputs, detecting drift and anomalies, and monitoring unexpected decisions
3. **Verifying identities:** Secure human users, APIs, autonomous agents, and machine identities
4. **Protecting the entire life cycle:** Training, fine-tuning, deployment, inference, and monitoring

The goal is not to trust a model simply because it resides inside the organization. Instead, every AI output should be considered a claim that must be continuously validated.

Q: What are the biggest challenges organizations face when implementing an AI governance framework such as ISO/IEC 42001?

A:

- ▶ **Competency gaps:** Organizations often lack professionals who understand AI governance, risk management, compliance, and life cycle management.
- ▶ **Shadow AI:** Employees increasingly use AI tools without governance, creating hidden risks and unknown data flows.
- ▶ **Life cycle governance:** Unlike traditional IT systems, AI systems continuously evolve through retraining and adaptation.
- ▶ **Explainability and transparency:** Organizations often struggle to understand why complex AI systems make specific decisions.
- ▶ **Third-party dependencies:** Many organizations use external models, cloud providers, and AI services they cannot fully control.

Q: With the myth of confidentiality on an AI model, can we say that the traditional data life cycle is not suitable here since we have risks such as reconstruction of sensitive data, or inability to fully delete data?

A: Not entirely, but it is no longer sufficient on its own. Traditional data life cycle concepts assume that data can be identified, stored, transferred, deleted, and controlled directly. AI fundamentally changes these assumptions. Once training data becomes embedded within a model, organizations may lose the ability to precisely identify, remove, or fully trace its influence.

Examples include model inversion attacks, training data reconstruction, memorization effects, leakage through outputs and retention beyond intended deletion schedules. Therefore, organizations should move from a classic “data life cycle” perspective toward an AI knowledge life cycle perspective that also considers training, fine-tuning, embeddings, model memory, inference outputs, and RAG knowledge stores.





Q: What are the most practical steps organizations can take to move from AI governance theory to effective implementation?

A: A practical roadmap would be:

- ▶ **Step 1 – Establish an AI inventory:** Identify where AI is actually being used. Shadow AI must become visible.
- ▶ **Step 2 – Define governance roles:** Introduce clear ownership across business, AI, IT, compliance, and risk functions.
- ▶ **Step 3 – Classify AI use cases:** Apply risk-based categorization consistent with the EU AI Act.
- ▶ **Step 4 – Implement AI-specific controls:** Include model validation, bias testing, monitoring, explainability controls, and data governance controls.
- ▶ **Step 5 – Build workforce competence:** Train auditors, risk officers, CISOs, business owners, and AI practitioners.
- ▶ **Step 6 – Establish continuous monitoring:** Governance must operate throughout the life cycle, not only during implementation.

Q: Is it possible to adapt the EU AI Act in Africa to help to manage AI?

A: Yes, although not necessarily as a direct legal copy. The EU AI Act principles are largely technology-neutral and can be adapted to African regulatory environments regardless of sector or country maturity. A localized approach could be to respect national legal systems, consider local economic realities, address sector-specific priorities, and maintain internationally recognized governance practices.

The combination of a localized regulatory framework with ISO/IEC 42001 provides a particularly attractive model because ISO/IEC 42001 is globally applicable and already aligns closely with many governance expectations reflected in the EU AI Act.

African countries do not necessarily need to copy the EU AI Act article by article. The greater value lies in adopting its governance philosophy: Risk-based oversight, accountability, transparency, human control, and continuous monitoring of AI systems.

Q: It is being said that CAIP is needed to complement ISO/IEC 42001 when implementing an AIMS. Can the ISO/IEC 42001 Lead Implementer lead the implementation of the AIMS and work with the subject-matter experts where required? Apart from people qualified in AI, you would need GDPR experts, legal experts, etc. What is your view on this?

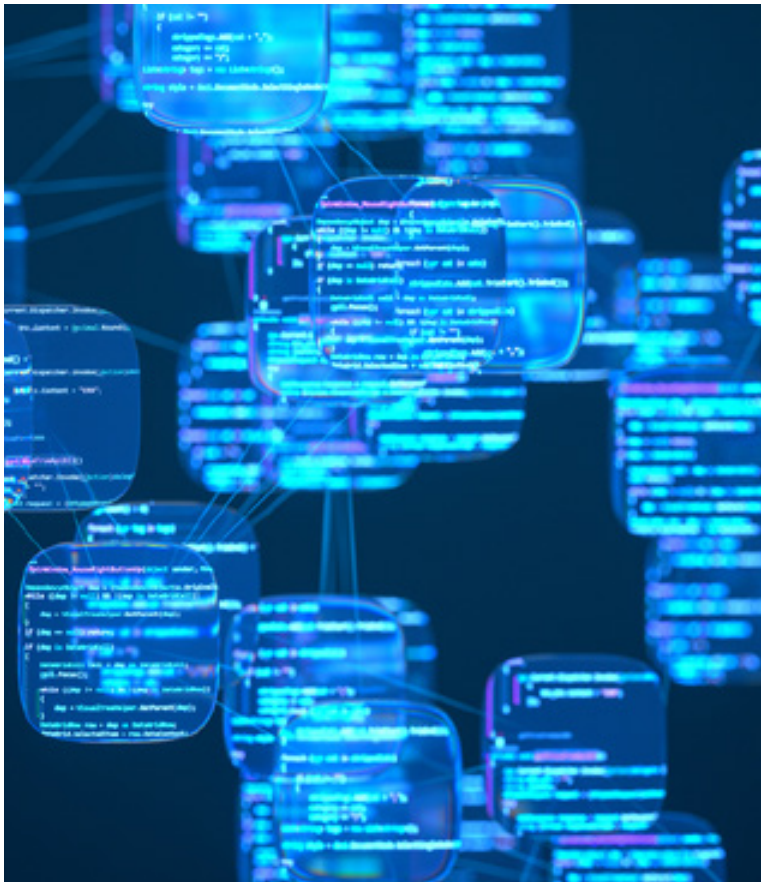
A: An ISO/IEC 42001 Lead Implementer can absolutely lead the implementation of an AIMS, just as an ISO 27001 Lead Implementer can lead an ISMS without personally being the foremost expert in cryptography, law, HR, or software development. The Lead Implementer's primary responsibility is to design, coordinate, and operationalize the management system.

However, AI governance is highly interdisciplinary. Effective implementation typically requires contributions from AI and data science specialists, legal and regulatory experts, GDPR/privacy specialists, information security professionals, business process owners and risk management functions.

The value of CAIP is therefore not that it replaces these specialists. Rather, it helps implementation leaders, auditors, CISOs, and governance professionals understand the AI-specific risks and governance concepts sufficiently well to ask the right questions, challenge assumptions, and coordinate the relevant stakeholders.

ISO/IEC 42001 Lead Implementer defines and implements the governance framework. Subject-matter experts provide specialist input. CAIP helps bridge the competence gap between governance and AI-specific realities.





Q: Who will be accountable for a deployed AI's actions? Is it the IT person who deployed it, their manager, the business leader who asked and funded it, or someone else?

A: Accountability can be delegated operationally, but it cannot be outsourced organizationally. The webinar repeatedly highlights that AI introduces accountability challenges because decisions emerge from a combination of model behavior, training data, operational deployment, business processes and human oversight. Therefore, accountability typically exists on several levels:

- ▶ **Technical accountability:** AI engineers, deployment teams, and model operators responsible for technical implementation and operation
- ▶ **Operational accountability:** Process owners and department managers responsible for how AI is embedded into business processes
- ▶ **Governance accountability:** Risk management, compliance and AI governance committees responsible for oversight and controls
- ▶ **Executive accountability:** The EU AI Act direction discussed during the webinar clearly points toward management-level accountability for AI decisions and governance.

The business leader who owns the business process ultimately remains accountable, even if deployment and operation are performed by technical teams.

Q: To remain relevant in ISMS, do I need both ISO 27001 and ISO/IEC 42001, or is ISO/IEC 42001 alone sufficient?

A: For information security professionals, I would strongly view these standards as complementary rather than alternative. ISO/IEC 42001 extends ISO 27001 rather than replaces it. AI governance introduces additional dimensions such as fairness, explainability, accountability, model governance, and life cycle monitoring. At the same time, ISO/IEC 42001 does not eliminate the need for access control, asset management, incident management, vulnerability management, security operations and classical information security governance.

Pragmatically, for the respective roles the recommended focus should be:

- ▶ **Security professional:** ISO 27001 and ISO/IEC 42001
- ▶ **AI governance professional:** ISO/IEC 42001 and AI expertise
- ▶ **Pure information security specialist:** ISO 27001 remains essential
- ▶ **Auditor of AI systems:** Both are highly valuable

ISO 27001 provides the foundation. ISO/IEC 42001 extends it into the AI domain. Neither completely replaces the other.

Q: How does ISO/IEC 42001 balance the need to encourage AI innovation while managing risks such as bias, lack of transparency, data privacy, and accountability?

A: A common misconception is that governance restricts innovation. In reality, ISO/IEC 42001 attempts to enable sustainable innovation by making it controllable. The framework does this through several mechanisms:

- ▶ **Risk-based governance:** Not every AI system receives the same level of governance. Controls are proportionate to risk and impact.
- ▶ **Life cycle management:** Instead of prohibiting innovation, organizations can experiment while ensuring monitoring, validation, and oversight throughout the life cycle.
- ▶ **Transparency and accountability:** Innovation remains possible as long as organizations can explain, govern, and monitor their systems appropriately.
- ▶ **Human oversight:** Rather than preventing automation, ISO/IEC 42001 seeks to ensure humans remain able to intervene where required.

Innovation without governance creates uncontrolled risk. Governance without innovation creates stagnation. ISO/IEC 42001 tries to create a structure in which both can coexist.

Q: What decision thresholds require human override? How do we quantify: benefit vs. exposure and productivity vs. risk?

A: There are no prescribed specific numerical thresholds. However, repeatedly stressed is the importance of human oversight, accountability, and escalation mechanisms. A practical governance approach is to require mandatory human intervention when decisions involve legal consequences, regulatory consequences, financial impact, safety impact, significant reputational impact, and irreversible actions. For quantification, organizations commonly define:

- ▶ **Benefit indicators:** Cost reduction, time savings, productivity increase, and quality improvement
- ▶ **Exposure indicators:** Financial loss potential, legal exposure, compliance impact, privacy impact, and operational disruption

The key principle should be: The greater the potential irreversible impact, the lower the threshold should be for human oversight.



Q: When organizations are already on the back foot, trying to gain buy-in to AIMS and mature, how can AI governance professionals adapt to increasing use of agentic AI and attitudes of innovation at all cost?

A: The challenge is that governance often moves slower than innovation. The webinar explicitly identifies this gap as a source of governance failure. Three adaptations are particularly important:

- ▶ **Move governance earlier:** Governance should be embedded into ideation and design phases rather than introduced after deployment.
- ▶ **Focus on guardrails, not roadblocks:** Governance teams should define boundaries, autonomy levels, escalation paths, and acceptable risk levels.
- ▶ **Prepare for agentic AI:** Organizations must clearly define what agents may do, which actions require approval, how actions are monitored, and who remains accountable.

Governance professionals should position themselves as enablers of safe innovation rather than opponents of innovation. That is the only scalable response to increasingly autonomous AI ecosystems.

Q: How do we adapt to AI auditing?

A: AI auditing requires auditors to shift from validating controls toward validating behavior. Future AI auditors should expand their competencies in:

- ▶ **AI governance:** ISO/IEC 42001, EU AI Act concepts, and life cycle governance
- ▶ **Data governance:** Lineage, representativeness, and quality assessment
- ▶ **AI security:** Prompt injection, model leakage, and poisoning attacks
- ▶ **Model assessment:** Explainability, robustness, and output validation
- ▶ **Continuous assurance:** Monitoring, scenario testing, and adversarial testing

The key transformation can be summarized as following:

- ▶ Control testing --> Behavior validation
- ▶ Point-in-time assessment --> Life cycle assessment
- ▶ Documentation review --> Scenario-based testing
- ▶ Infrastructure focus --> Data, model, and outcome focus

AI is a broader discipline than IT auditing, requiring interdisciplinary competencies, continuous assessment, and understanding of how AI systems actually behave in practice.

Q: About AI hallucinations and misleading output, how do you evaluate if AI is hallucinating?

A: One of the key challenges highlighted in the webinar is that hallucinations often appear plausible and convincing, making them difficult to detect automatically. A practical evaluation approach combines several methods:

- ▶ **Ground truth verification:** Compare AI output against authoritative source documents, validated databases and subject matter expert knowledge. If the output cannot be traced back to known facts, additional scrutiny is required.
- ▶ **Consistency testing:** Ask the model the same question in different ways, follow-up questions, requests for sources and reasoning. Large inconsistencies may indicate unreliable output.
- ▶ **Explainability assessment:** Evaluate whether conclusions are supported by evidence, reasoning is internally coherent, assumptions are transparent. The webinar highlights explainability as a major governance and audit criterion.
- ▶ **Scenario-based testing:** Auditors increasingly test prompts, adversarial scenarios, edge cases to observe actual AI behavior under realistic conditions.
- ▶ **Human review for critical decisions:** For high-impact decisions, hallucination detection should never rely solely on automation. Human validation remains essential.

Hallucinations cannot be completely eliminated. Governance therefore focuses on detecting, monitoring, and mitigating them rather than assuming they will never occur.

Q: ISO/IEC 42001 is largely written around a single AI system. In a multi-agent setup where authority is delegated across several autonomous agents, accountability gets distributed. Do the Annex A controls cover that, or do you need a separate agent responsibility/accountability matrix layered on top of the AIMS?

A: ISO/IEC 42001 provides governance principles around accountability, oversight, role definition, human intervention, life cycle governance. However, the presentation repeatedly emphasizes that accountability becomes increasingly complex in distributed AI ecosystems.

An agent-specific responsibility and accountability matrix is highly advisable as an implementation mechanism within the AIMS. Such a matrix could define agent ownership, delegated authority, escalation paths, human override responsibilities, monitoring responsibilities, and incident response ownership.

While there is no explicitly prescribed such a matrix, the governance challenges it highlights strongly suggest the need for more granular accountability structures in multi-agent environments.

Q: Why do calls for oversight and supervision as part of operations face outright rejection, despite the logic of the idea?

A: In many organizations, governance is perceived as slowing innovation, reducing agility, creating additional workload, and challenging existing autonomy. Several psychological and organizational factors often contribute, including perceived loss of efficiency, automation bias, illusion of control, innovation culture, and lack of visible incidents.

The most effective response is a risk-reduction and business-enablement mechanism. The objective of governance is not to stop innovation. The objective is to ensure that organizations remain in control as AI systems become more autonomous, scalable, and difficult to understand.

Q: Can AI tools not be SOC-certified to ensure confidentiality?

A: SOC reports can be extremely valuable, but they do not fully solve the confidentiality problem described in the webinar. A SOC report may provide assurance regarding organizational controls, security operations, access management, infrastructure protection, or operational processes.

However, the webinar's "myth of confidentiality" discussion focuses on different challenges like data embedded into model behavior, model inversion attacks, unintended disclosure, inability to completely remove training influences, and prompt-based leakage. A SOC-certified AI provider may demonstrate strong organizational controls, but this does not automatically eliminate model-specific confidentiality risks.





Q: Given the pace of AI technology, AI work volume and number of new AI developments/applications that are not adapted to an agreed standard (like ISO/IEC 42001 etc.), how can a unified standard be implemented worldwide?

A: A fully unified global regulatory model faces obvious challenges, like different legal systems, different economic priorities, different political approaches, and different levels of technological maturity. This suggests that AI governance will likely evolve through several layers:

- ▶ **International frameworks:** Examples include ISO standards, global governance principles, and industry best practices.
- ▶ **Regional regulation:** Examples include EU AI Act, national AI legislation, and sector-specific requirements.
- ▶ **Organizational governance:** Organizations must govern their own use of AI regardless of external regulation.

Q: There is a rising sense of urgency among employees and pressure from organizations to leverage AI capabilities. In most cases, this drives the use of free tools, leading to the upload of sensitive company data into these platforms. How can we mitigate this urgency while balancing corporate budgets against the cost of premium AI subscriptions?

A: A purely prohibitive approach is unlikely to succeed. Employees typically turn to free tools because they need productivity gains. Approved alternatives are unavailable and approval processes are too slow. Combine governance, awareness, and practical alternatives.

1. **Provide safe alternatives:** The easiest way to reduce Shadow AI is to provide an approved alternative that satisfies business needs.
2. **Classify data:** Not all information requires the same protection level. Organizations should clearly define what may be entered into public AI tools, what may only be entered into approved corporate solutions and what must never leave the organization.
3. **Build awareness:** Many employees do not intentionally violate policy; they simply underestimate AI-related risks.
4. **Focus premium licensing on high-impact users:** Instead of purchasing enterprise licenses for everyone, organizations can often achieve substantial risk reduction by prioritizing knowledge workers, research functions, software engineering teams and high-volume AI users.
5. **Measure the cost of leakage:** A single disclosure of source code, customer data, research results and strategic documents may cost far more than a controlled enterprise AI deployment.



Q: What practical experience or technical experience should a legal professional have in order to become AI officer?

A: I repeatedly emphasized that AI governance is interdisciplinary and requires expertise beyond legal compliance alone. An AI Officer does not necessarily need to become a data scientist, but should understand how AI systems are designed, trained, deployed, monitored, and governed throughout their life cycle. For a legal professional, the most valuable additional competencies would be:

- ▶ Technical understanding: AI life cycle fundamentals (training, deployment, retraining, and monitoring), AI-specific risks (bias, hallucinations, explainability issues, and model drift), AI supply-chain risks and third-party dependencies, and understanding of LLMs, RAG architectures, and AI services at a conceptual level
- ▶ Governance and risk management: AI risk assessment methodologies, ISO/IEC 42001 governance principles, accountability and oversight structures, and auditability and documentation requirements
- ▶ Regulatory expertise: GDPR, EU AI Act, sector-specific regulations, and contractual and liability considerations.

An effective AI officer is typically not purely a lawyer, purely a technologist, nor purely a risk manager. The role increasingly requires the ability to translate between all three domains.

Q: What actions can be taken to stop the use of 123 Notetaker AI on Microsoft Teams which can be downloaded that may lead to sensitive information available to unauthorized individuals or organizations.

A: The discussion on Shadow AI is directly relevant here. Uncontrolled AI tools create governance blind spots, hidden data flows, and confidentiality risks. Potential governance measures include:

1. **Define an AI usage policy:** Clearly specify approved AI notetaker solutions, prohibited solutions, permitted meeting types and handling of confidential information.
2. **Establish an approval process:** Require review and approval before any AI meeting assistant is used.
3. **Classification-based restrictions:** Certain meetings may prohibit recording, transcription, AI summaries, based on information sensitivity.
4. **Awareness and training:** Many users do not realize where recordings are stored, who can access transcripts, whether data is processed externally.
5. **Vendor risk assessment:** Third-party AI services are a significant governance challenge. AI notetakers should therefore be assessed similarly to any other supplier.

Organizations cannot govern AI usage they cannot see. Visibility and clear ownership are the first controls.





Q: How can new AI standards assist in mitigating or guarding against anthropomorphism?

A: Anthropomorphism refers to the tendency of users to attribute human characteristics, judgment, emotions, or understanding to AI systems. Standards such as ISO/IEC 42001 can help by requiring:

- ▶ **Transparency:** Users understand when they are interacting with AI rather than a human decision-maker.
- ▶ **Human oversight:** Important decisions remain subject to human review.
- ▶ **Explainability requirements:** Organizations document system limitations and intended use cases.
- ▶ **Governance of automated decisions:** Decision authority and accountability remain assigned to humans.

AI should be trusted appropriately, not treated as either infallible or human. Effective governance helps maintain that balance.

Q: What governance can be placed on AI providers on how they handle customer data if there needs to be stringent regulations on how they are using customer data and accountability if data is mishandled by AI?

A: Organizations must govern external AI services through contractual, technical, and governance mechanisms. Important controls include:

- ▶ **Transparency requirements:** Providers should explain how customer data is processed, where it is stored, whether it is used for training, and applicable retention periods.
- ▶ **Contractual controls:** Contracts should address confidentiality obligations, data ownership, audit rights, and incident notification requirements.
- ▶ **Supply chain governance:** AI providers should be integrated into vendor assessment processes, risk management activities and continuous monitoring programs.
- ▶ **Accountability:** Accountability remains with the deploying organization, even when external providers are involved.

Q: How relevant is it to put AI risk-related events to the board for discussion?

A: Increasingly, it is highly relevant. Accountability becomes a management-level responsibility. AI failures can have legal, financial, operational, and reputational consequences, and AI governance is becoming a regulatory obligation.

Board-level discussion is particularly appropriate when risks involve regulatory exposure, significant financial impact, strategic business processes, customer trust, reputation or high-risk AI systems. Not every AI incident requires board involvement. However, material AI risks should increasingly be treated like major cybersecurity, compliance, or enterprise risk events.

Q: How does ISO/IEC 42001 help organizations identify and manage AI-related risks before they impact operations or stakeholders?

A: The core purpose of ISO/IEC 42001 is proactive governance rather than reactive response. The standard introduces a life cycle-oriented approach to risk management. There are several mechanisms available:

- ▶ **Risk identification:** Organizations identify risks across training data, models, deployment, third-party providers, and operational use.
- ▶ **AI-specific controls:** Controls include model validation, data quality controls, performance monitoring, and oversight mechanisms.
- ▶ **Continuous monitoring:** Governance continues after deployment through monitoring, incident management, and performance assessment.
- ▶ **Accountability structures:** Clear responsibility assignment ensures risks are addressed before incidents occur.

ISO/IEC 42001 helps move organizations from reacting to AI failures toward systematically preventing and detecting risks throughout the AI life cycle.

Q: There are different regulatory frameworks put in place that companies need to follow, and if they do not comply, they are fined a certain amount. Does that apply in the AI space?

A: AI regulation is moving from voluntary guidance toward enforceable obligations. The slides specifically reference EU AI Act requirements, enforcement mechanisms, liability considerations, and regulatory exposure.

Therefore, today certain AI-related activities may already trigger GDPR penalties, sector-specific regulatory actions or contractual liability, depending on context.

Going forward, the EU AI Act introduces further regulatory requirements and enforcement mechanisms for certain AI systems. We explicitly highlighted accountability, auditability, risk management, and governance obligations.

AI governance is progressively moving from a best-practice discussion to a compliance and regulatory requirement.

Q: What are the key controls that an Information Systems Manager should put in place to ensure effective AI governance?

A: The webinar provides numerous examples of controls that can be translated into a practical governance baseline.

- ▶ **Governance controls:** AI policy framework, AI inventory, ownership, accountability structures, and AI risk assessments
- ▶ **Data governance controls:** Data classification, data quality management, lineage tracking and handling of sensitive information
- ▶ **AI life cycle controls:** Model validation, approval before deployment, retraining governance, and performance monitoring
- ▶ **Security controls:** Prompt protection, access management, monitoring of AI services, and supply-chain risk management
- ▶ **Human oversight controls:** Escalation procedures, review mechanisms, accountability frameworks, and training and awareness
- ▶ **Monitoring controls:** Incident management, continuous assurance, audit programs, and metrics and reporting

The most effective AI governance combines technical controls, operational oversight, accountability structures, and continuous monitoring. Governance must cover the entire AI life cycle.

Q: Most business today is conducted cross-border. Apart from the EU AI Act, what are the other regulators doing (e.g. Americas, Asia, Africa)?

A: From the governance perspective discussed during the webinar, several common themes are emerging globally, like risk-based governance, accountability for AI decisions, transparency and auditability, human oversight, third-party and supply-chain governance, or documentation and traceability requirements.

The practical challenge for multinational organizations is therefore less about complying with a single regulation and more about establishing a governance framework that can accommodate different regulatory expectations across regions. For that reason, ISO/IEC 42001 can be particularly valuable because it provides a governance framework that is not tied to a single jurisdiction and can complement region-specific legal requirements.

Q: Why do we use AI if it has evidence of challenges?

A: AI is simultaneously a source of opportunity and risk. Organizations continue to adopt AI because it can provide increased productivity, faster analysis, better scalability, improved automation, enhanced decision support, and advanced cybersecurity capabilities. AI is both a force multiplier for cyber threats, and an enabler for more advanced defense capabilities.

The existence of risk does not automatically justify avoiding a technology. Organizations routinely use technologies that introduce risk, provided those risks can be governed and managed. The challenge therefore becomes whether it can be governed responsibly. The overall conclusion is that AI adoption is likely unavoidable in many industries. The real differentiator will be governance maturity rather than AI adoption itself.



Q: How do we practically classify AI systems into risk levels, and who approves that classification?

A: The EU AI Act introduces a risk-based classification model that determines the level of governance and regulatory scrutiny required. A practical approach typically considers:

- ▶ **Risk factors:** Impact on individuals, business criticality, degree of automation, potential safety implications, legal or regulatory consequences, and data sensitivity.
- ▶ **Governance responsibility:** Classification should usually be performed jointly by business owners, risk management, compliance, AI governance functions, and technical stakeholders.
- ▶ **Approval:** Classification approval should generally occur through an established governance process with designated ownership and oversight.

Q: If we include AI into governance systems, who will guarantee that AI will not cause harm to humans?

A: No one can provide such a guarantee. AI can fail, hallucinate, and AI decisions can have unintended consequences. Governance must assume failure rather than perfection. The purpose of governance is therefore not to eliminate all risk.

Governance attempts to reduce risk, detect failures, limit impact, establish accountability, and maintain human oversight. This is similar to aviation, medicine, or cybersecurity. Governance frameworks improve safety but cannot provide absolute guarantees.

Q: Would you recommend CAIP certification for technical staff, for example a veterinary technician, as a support as we move towards AI adoption in the future or management by the IT department suffices?

A: AI governance is not solely an IT responsibility. AI increasingly affects operations, compliance, healthcare, research, business processes, and decision-making. For technical professionals, such as veterinary technicians, the decision should depend on their anticipated involvement in AI-enabled processes. CAIP may be valuable if they will actively use AI systems, they participate in AI-supported diagnostics, they supervise AI-assisted workflows, they are expected to understand AI-related risks and limitations.

CAIP may be less critical if their involvement is minimal or AI governance remains entirely centralized.

However, trustworthy AI adoption depends on broader organizational competence. AI becomes embedded into operational processes; basic AI governance knowledge will increasingly be beneficial beyond IT functions alone.

Q: How do we detect manipulated outputs?

A:

- ▶ **Behavior validation:** Compare outputs against trusted sources, known business rules and expected behaviors.
- ▶ **Scenario testing:** Auditors increasingly use prompt testing, adversarial testing, and simulation scenarios to identify unusual model behaviors.
- ▶ **Continuous monitoring:** I emphasized the requirement for monitoring outputs rather than relying solely on documentation reviews.
- ▶ **Human oversight:** For high-impact decisions, independent review remains an essential control.

AI systems may be manipulated without appearing technically compromised. Governance must therefore focus on detecting abnormal behavior, not only system failures.



Q: “AI Governance frameworks such as ISO/IEC 42001 alone are insufficient without qualified professionals.” Discuss this statement critically, highlighting the role of AI governance knowledge, risk management, compliance, auditing, and certification.

A: This statement is largely consistent with one of the webinar’s central messages: “Frameworks define what must be done; people ensure it is actually done”. ISO/IEC 42001 provides governance structures, life cycle processes, accountability mechanisms and risk-management requirements.

Frameworks cannot identify every AI risk automatically, interpret complex business contexts, challenge assumptions, investigate failures, or exercise professional judgment. Qualified professionals therefore remain critical in risk management, governance implementation, auditing, regulatory compliance, and oversight activities. At the same time, expertise without governance frameworks can also be problematic. Highly skilled professionals operating without defined controls, accountability structures and governance processes may create inconsistent outcomes. Therefore, neither frameworks nor professionals alone are sufficient. Real AI governance emerges from the combination of governance structures and qualified people.

Q: What contingencies will be there in case AI audit shows an agent or a model or data-based risk. How do we manage uninterrupted IS operations?

- A:** A practical response strategy includes:
- ▶ **Risk-based remediation:** Not every finding requires immediate shutdown. Controls should reflect severity, likelihood, and operational impact.
 - ▶ **Escalation processes:** I recommend structured incident management, reporting, and escalation mechanisms.
 - ▶ **Human override:** Critical decisions may require increased review, temporary manual controls, and additional approvals.
 - ▶ **Fallback mechanisms:** Organizations should maintain non-AI alternatives, manual procedures, and contingency workflows.
 - ▶ **Continuous monitoring:** Use monitoring to verify whether remediation is effective.



Q: How do we detect and respond to prompt injection or model compromise in real time? Where in the life cycle are these risks assessed and mitigated?

A: For detection, monitoring AI behavior, prompt testing, adversarial testing, continuous oversight, and output validation are recommended. For response, potential governance actions like blocking suspicious interactions, escalating incidents, applying additional review and updating safeguards and controls are recommended.

From the life cycle perspective risks must be addressed across the entire life cycle:

- ▶ **Training stage:** Data validation and protection from poisoning attacks
- ▶ **Validation stage:** Model testing and adversarial assessments
- ▶ **Deployment stage:** Secure architectures, controlled integration, and interface security
- ▶ **Operational stage:** Continuous monitoring, incident management, and behavior verification

Prompt injection and model compromise are not purely technical problems. They require continuous governance, monitoring, oversight, and life cycle-based risk management.

Q: What is the role of AI professionals in EU AI Act readiness?

A: Compliance with the EU AI Act requires people who understand AI systems, governance, risk management, documentation, and auditability. AI professionals help organizations translate regulatory requirements into operational controls and governance processes. Their role typically includes governance implementation, risk management, documentation and auditability, and regulatory readiness. AI professionals bridge the gap between regulatory requirements and practical implementation. The EU AI Act defines obligations; qualified professionals help organizations operationalize them.



Q: What are the new trust boundaries in AI ecosystems?

A: Traditional trust boundaries were largely defined by networks and infrastructure. In AI environments, those boundaries shift significantly. Trust boundaries move from networks toward data, models, services, and automated decision-making chains.

New trust boundaries could include:

- ▶ **Data flows:** Organizations must understand where data originates, how it is processed, and where it is transmitted.
- ▶ **AI models:** “Trust considerations” now include model behavior, training origins and validation status.
- ▶ **AI services and providers:** External AI services may operate outside the organization’s direct control.
- ▶ **AI agents:** Autonomous or semi-autonomous systems create new trust relationships requiring governance and oversight.
- ▶ **AI pipelines:** Models, APIs, retrieval mechanisms, and integrations create distributed trust boundaries across the AI ecosystem.

Q: Why is identity and access management (IAM) important for AI systems?

A: AI introduces entirely new categories of identities. These identities may include human users, APIs, AI agents, machine-to-machine services, and automated workflows. Strong IAM helps organizations control access, reduce risk by preventing unauthorized access to sensitive prompts, training data, and AI services, support accountability by tracing actions to responsible identities, and enable governance through monitoring, auditability, and oversight.

In AI environments, identity management applies not only to people but increasingly to machines, APIs, services, and autonomous agents.

Q: What is meant by data protection across the AI life cycle?

A: Data protection must extend beyond traditional concepts of “at rest” and “in transit.” Data protection must be applied throughout all phases of the AI life cycle:

- ▶ **Data collection:** Ensure appropriate and lawful acquisition of data.
- ▶ **Training:** Protect training data against unauthorized access, manipulation and poisoning attacks.

- ▶ **Deployment:** Control access to operational AI systems.
- ▶ **Inference:** Prevent leakage of sensitive information through prompts and outputs.
- ▶ **Monitoring and maintenance:** Continuously assess confidentiality, integrity, and security throughout model operation.

Data protection in AI environments is a life cycle activity rather than a one-time security control.

Q: Where does AI store its data outputs (like Copilot or Grok)?

A: We discussed provider transparency, confidentiality challenges, and third party AI services, but we cannot provide any specific information about where particular commercial products such as Microsoft Copilot or Grok store prompts, outputs, or conversation history, as it is the matter of the development details of each and every commercial product.

Organizations should understand how providers process data, where data is handled, what transparency exists, what contractual and governance controls are in place, and what risks arise from third-party AI services. The details of the implementations and storage location can be obtained from the respective commercial product owners.

Before using any external AI service, organizations should verify the provider’s documented data handling, storage, retention, and governance practices rather than assuming that outputs are stored in a particular location.





Dr. Roman Krepki

Senior Manager for
Risk at Forvis Mazars

Dr. Roman Krepki is a senior consultant at Forvis Mazars GmbH & Co. KG in Stuttgart, Germany, where he advises organizations on information security, digital resilience, governance, and risk management with a particular focus on the financial sector. His expertise includes supporting clients in the implementation of ISO/IEC 42001, DORA, and NIS2, as well as conducting risk assessments, and continuous improvement of security systems as well as establishing security frameworks for the secure and resilient use of Artificial Intelligence. In addition, he is a frequent trainer, and speaker on information security, digital resilience, regulatory compliance, and Artificial Intelligence.



Dr. Dr. h.c. Christian Krepki

Senior Manager for Cybersecurity and
Data Protection at Forvis Mazars GmbH & Co. KG

Senior manager at Forvis Mazars
Stuttgart, Germany. He advises
on security, data protection,
finance, risk, and compliance,
ISO/IEC 27001, TISAX®,
and the EU AI Act. His work
includes the design, implementation,
improvement of management
ensuring effective governance
and responsible use of Artificial
Intelligence. He regularly serves as an auditor,
information security, cybersecurity,
and AI governance.



eLearning

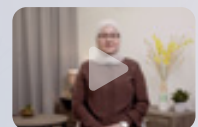
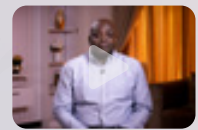
eLearning

eLearning

eLearning



ISO/IEC 27001 Foundation (French) – Updated



Strengthen Your Expertise in Information Security and Privacy

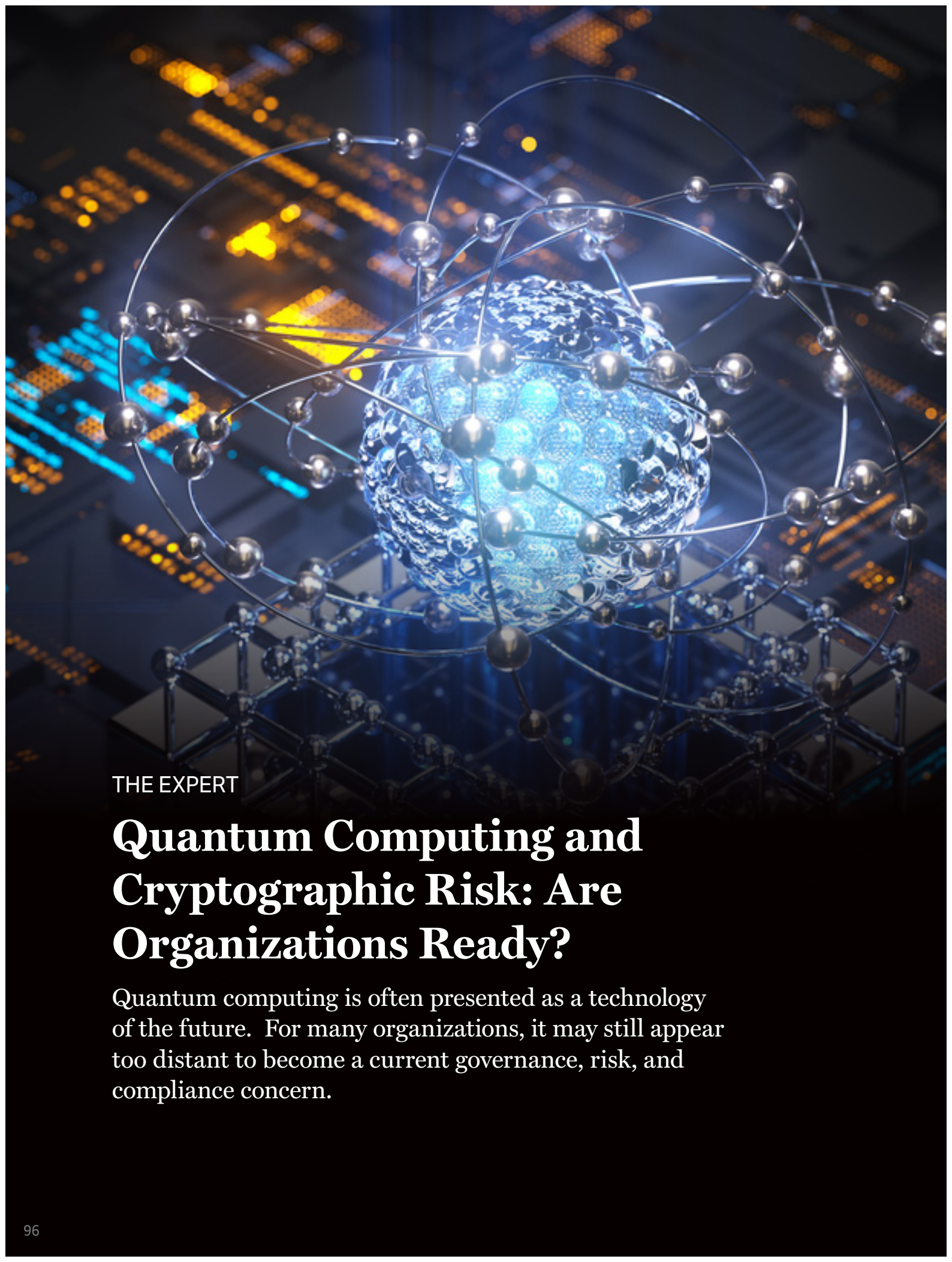
Stay ahead of evolving security and privacy requirements with PECB's flexible eLearning training courses.

Explore the new and updated training courses in the ISO/IEC 27001 and ISO/IEC 27701 families:

- [ISO/IEC 27001 Foundation \(French\)](#) – Updated
- [ISO/IEC 27701 Foundation \(English\)](#) – New
- [ISO/IEC 27701 Lead Implementer \(English\)](#) – New
- [ISO/IEC 27701 Lead Auditor \(English\)](#) – New

Learn at your own pace, wherever you are, with expert-led, interactive content and flexible access designed to fit your professional schedule. Build knowledge, strengthen your credentials, and advance your career in information security and privacy.

[LEARN MORE →](#)



THE EXPERT

Quantum Computing and Cryptographic Risk: Are Organizations Ready?

Quantum computing is often presented as a technology of the future. For many organizations, it may still appear too distant to become a current governance, risk, and compliance concern.

For many organizations, it may still appear too distant to become a current governance, risk, and compliance concern. However, its possible impact on cryptography makes the question more urgent than it first seems.

Today, organizations depend on cryptographic mechanisms to protect sensitive data, secure online communications, authenticate users, and verify digital signatures. However, the development of quantum computing over recent years is revealing the future limitations of some widely used public-key cryptographic mechanisms, especially RSA and elliptic-curve cryptography (ECC). A powerful quantum computer could solve the mathematical problems on which the security of these mechanisms depends.

Organizations have an interest in considering quantum risk today. Information encrypted now may still be sensitive when quantum capabilities become available. Furthermore, replacing cryptographic mechanisms across complex information systems may require several years. Waiting for the threat to become operational could therefore leave organizations with insufficient time to respond. This situation creates a difficult question for organizations: how can they prepare for a risk when the threat is not yet operational and its timeline remains uncertain? The response will require the involvement of different actors, from cybersecurity and IT teams to risk managers, business owners and suppliers. GRC can provide the structure needed to understand the exposure, assess the risk, and coordinate the transition.

This article examines how quantum computing may affect current cryptographic mechanisms, why organizations should prepare today, the challenges it creates for governance, risk, and compliance, and the main steps toward quantum readiness.

This topic will also be explored at the upcoming [PECB Conference](#) in Rome, where a dedicated panel discussion will examine quantum computing and its implications for organizations.

What Quantum Computing Changes in Cryptography

The main objective of encryption is to transform readable information into an unreadable form in order to prevent unauthorized access. And the information can be recovered only by using the appropriate cryptographic key. In today's computing environment, organizations mainly rely on two types of cryptographic mechanisms: **symmetric** and **asymmetric cryptography**.

Symmetric encryption uses the same secret key to encrypt and decrypt information, with AES being a common example. Asymmetric cryptography uses two mathematically linked keys: a public key that can be shared and a private key that must remain secret. **RSA** and **ECC** are widely used asymmetric mechanism.

These mechanisms are considered secure against classical computers when they are correctly implemented and use appropriate key sizes. Their security depends on mathematical problems that are easy to perform in one direction but extremely difficult to reverse. For example, RSA relies on the difficulty of finding the original prime numbers used to produce a very large number. When an appropriate key size is used, even a powerful classical computer would need an extremely long time to find these numbers. This mathematical difficulty is what makes RSA secure today.

However, progress in quantum computing could change this situation. Quantum computers process information differently from classical computers. A classical computer uses bits, and each bit has one value at a time, either 0 or 1. A quantum computer uses quantum bits, called qubits. Before it is measured, a qubit can exist in a combination of the two reference states, called a superposition. Quantum algorithms manipulate these states and use interference to reduce incorrect possibilities and increase the probability of obtaining a useful result. This does not make quantum computers faster for every task, but it can give them an important advantage when solving some specific mathematical problems.

The main cryptographic concern comes from [Shor's algorithm](#). A sufficiently powerful quantum computer could use it to solve the mathematical problems protecting RSA and ECC. This could compromise key establishment and digital signatures, affecting confidentiality, integrity, and authenticity. However, quantum computing does not affect every cryptographic mechanism in the same way. AES, which is a symmetric encryption algorithm, is not broken by Shor's algorithm. The main concern is the future vulnerability of public-key cryptography on which many digital systems currently depend.

Why Organizations Should Prepare Today?

As mentioned above, a powerful quantum computer capable of breaking current public-key cryptography is not available today, and no one can predict exactly when it will become operational. However, this uncertainty does not mean that organizations should wait before assessing the risk. The need to prepare is mainly justified by the lifetime of sensitive information and the time required to change cryptographic systems.

The first concern is known as HNDL, or “Harvest Now, Decrypt Later.” It means that an attacker can collect encrypted information today and store it for several years, until sufficiently quantum capabilities become available and allow the attacker to decrypt it.

The level of risk depends largely on the lifetime of sensitive information and how long it must remain confidential. This may concern government information, personal data, research, intellectual property, or strategic business information. **The important question is not only how information is protected today, but also how long it needs to remain confidential.**

A second concern relates to digital signatures. RSA and ECC are used not only to establish secure communications, but also to verify the identity of a signer and confirm that information has not been modified. If these mechanisms become vulnerable, attackers could potentially forge signatures, impersonate trusted parties, or distribute malicious software as a legitimate update. This creates risks to integrity and authenticity.

Finally, changing cryptographic mechanisms is not as similar as installing a simple software update. Cryptography is present in applications, certificates, protocols, devices, and services provided by external suppliers. Some old systems may not support new algorithms, and suppliers may not all be ready at the same time. Organizations must first identify these dependencies, assess their importance, test new solutions, and manage compatibility. **This process can require several years. That is why organizations should start preparing today.** This does not mean that every system must be replaced immediately. It simply means giving the organization enough time to understand the risk, define its priorities and manage the transition step by step.

Quantum Computing Risk as a GRC Challenge

The different concerns discussed above show that the risks quantum computing creates for current cryptographic mechanisms are not limited to technical issues for IT and cybersecurity teams. They also represent a challenge for governance, risk, and compliance.

From a governance perspective, organizations need to define who is responsible for monitoring this risk and preparing the response. The subject can involve different functions, including information security, IT, risk management, legal, procurement, and business departments. Without clear responsibilities and coordination, each function may consider that the subject belongs to another department. Management should, therefore, ensure that quantum risk is considered in future security and technology decisions.

The level of risk is not the same for every organization. It depends on the type of information, how long it must remain protected, the cryptographic mechanisms used and the time needed to change them. For example, information that must remain confidential for 15 years is more exposed than information that loses its value after a short period. Each organization should, therefore, assess its own situation and give priority to sensitive information and critical systems.

Quantum risk may also create compliance concerns. Organizations have legal, regulatory, or contractual obligations to protect personal, financial, or strategic information. If the cryptographic protection used today may become vulnerable while the information still needs to remain confidential, this risk should be considered as part of compliance management.

Another important concern relates to external suppliers. Organizations depend on cloud services, software, certificates, equipment, and other services provided by third parties. Their ability to change cryptographic mechanisms will also depend on the preparation and migration plans of these suppliers. For this reason, supplier monitoring and future contractual requirements should be included in the preparation.

Organizations do not need to create a completely new risk-management framework to manage this subject. [ISO/IEC 27005](#) can help them identify, analyze, and evaluate scenarios related to long-term confidentiality, digital signatures, and migration difficulties. [ISO/IEC 27001](#) can support the governance, treatment, and monitoring of these risks through the existing information security management system.

The objective is not to predict the exact date when a powerful quantum computer will become available. It is to understand the organization’s exposure and start making the necessary decisions before the situation becomes urgent.





From Awareness to Quantum Readiness

Understanding the risk is a first step, but it does not mean that an organization is ready. Quantum readiness does not require the immediate replacement of all existing cryptographic mechanisms. It means that the organization knows its exposure, defines its priorities, and prepares a controlled transition.

The first action is to identify the information that needs to remain confidential for a long period. Organizations should determine where this information is stored, how it is exchanged and which cryptographic mechanisms protect it. This assessment should begin with sensitive information and critical systems instead of trying to analyze everything at the same time.

Organizations also need to identify where RSA and ECC are used. These mechanisms may be present in applications, digital certificates, communication protocols, digital signatures, devices and services provided by external suppliers.

Building this inventory may be difficult because some cryptographic mechanisms are hidden inside applications or are not clearly documented.

One of the main solutions to this risk is Post-Quantum Cryptography, known as PQC. It refers to cryptographic algorithms designed to resist attacks from both classical and quantum computers. These algorithms can run on current computers, so organizations do not need to have a quantum computer to use them.

In 2024, NIST published its first three final PQC standards. ML-KEM is mainly used to establish shared secrets, while ML-DSA and SLH-DSA are used for digital signatures. These standards give organizations an important direction for preparing their future [migration](#).

Nevertheless, adopting PQC is not the only action required. Organizations still need to identify their cryptographic dependencies, assess the related risks, test the new solutions and manage their compatibility with existing systems. New algorithms may also have different requirements in terms of performance, key sizes, and signatures. External suppliers are also an important part of the preparation. An organization may be ready to change its own systems but still depend on software, cloud services, or equipment that does not yet support post-quantum solutions. Organizations should, therefore, ask their critical suppliers about their migration plans, timelines, and future compatibility.

Another important element is crypto-agility. It means the ability to replace cryptographic algorithms, keys, or certificates without rebuilding the entire system or causing serious disruption. This is important because technologies, standards, and threats may continue to change even after the first migration.

Finally, organizations should [prepare a gradual migration](#) plan that defines priorities, responsibilities, budgets, and testing needs. A quantum-ready organization is not one that has already replaced every vulnerable mechanism. It is an organization that understands where the risk exists and prepares its transition step by step.

Conclusion

Quantum computing is still developing, and a computer capable of breaking current public-key cryptography is not available today. However, this does not mean that quantum risk should be ignored until such a computer becomes operational. Information collected today may still be sensitive in the future, while changing cryptographic systems can require several years.

Organizations should begin by understanding their exposure, identifying their priorities, and preparing a controlled migration. Existing risk-management frameworks can support this preparation, while PQC standards provide an important direction for the transition. Existing risk-management frameworks can support this preparation, while PQC standards provide an important direction for the future migration.

Are organizations ready? For many of them, the answer is probably not yet. Cryptographic dependencies are not always known, supplier plans may remain unclear, and quantum risk may still be considered a distant technical subject. However, becoming ready is a gradual process. It starts by recognizing the risk, assigning responsibilities and preparing the transition step by step.

The future of quantum computing remains uncertain, but preparing for its risks is already a shared responsibility across the organization.



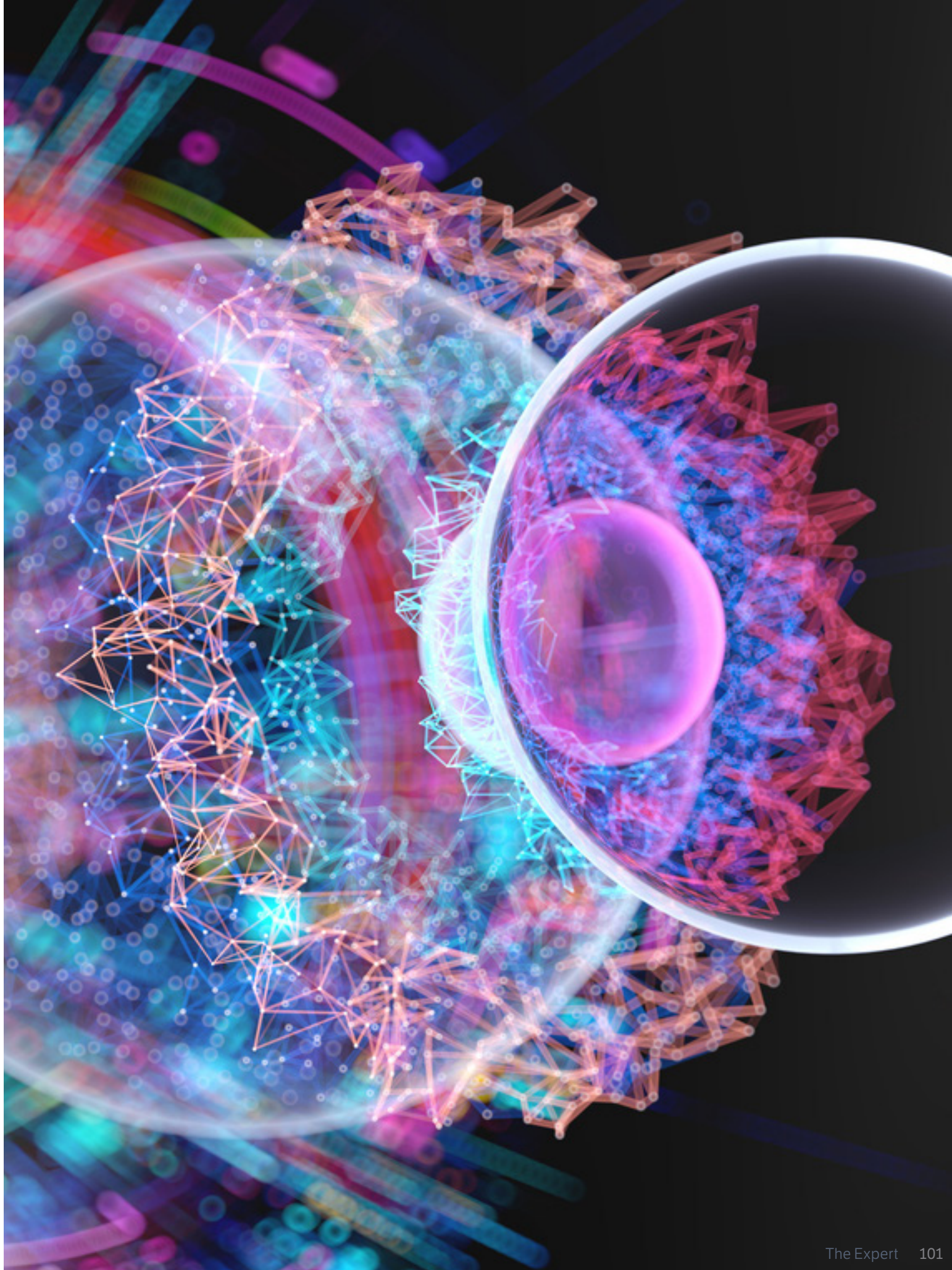
Hamza Matine

Internal Auditor and Risk
Manager

Hamza Matine holds a PhD and serves as an Internal Auditor and Risk Manager with deep expertise in information security and risk governance. He is certified as an ISO/IEC 27001 Lead Implementer and Lead Auditor, ISO/IEC 27005 Risk Manager, and EBIOS Risk Manager, bringing both strategic and technical dimensions to his practice.

With hands-on experience in designing and auditing risk management frameworks, Hamza combines academic rigor with real-world GRC expertise, making him a trusted voice in the fields of information security, audit, and organizational resilience.







WORK-LIFE BALANCE

The Executive Resilience Imperative: Work-Life Balance in an Era of Constant Risk

The modern executive no longer operates in a
predictable business environment.

Today's leaders navigate an increasingly interconnected world shaped by geopolitical instability, cyber threats, artificial intelligence, economic uncertainty, digital transformation, and an always-on information cycle. For executives leading organizations in technology, security, governance, defense, and critical infrastructure sectors, the stakes are even higher.

A decade ago, work-life balance was often viewed as a personal wellness topic, something discussed in leadership retreats or human resources seminars. Today, it has become a strategic leadership issue directly connected to resilience, decision-making quality, organizational culture, talent retention, operational continuity, and long-term business sustainability.

Executives are now expected to lead through disruption while simultaneously anticipating future risks that evolve faster than many governance models can adapt. Boards demand agility. Markets demand innovation. Employees demand transparency and flexibility. Regulators demand accountability. Meanwhile, cyber-attacks, disinformation campaigns, AI-related risks, and global instability continue to place unprecedented cognitive pressure on leadership teams. In this environment, the traditional idea of “work-life balance” is no longer sufficient. The future belongs to leaders who understand executive resilience.

The organizations that will thrive in the coming decade will not necessarily be those with the most advanced technologies alone, but those led by individuals capable of sustaining clarity, adaptability, ethical judgment, and strategic focus under constant pressure.

The Hidden Cost of Executive Overload

Many executives operate under a dangerous assumption: high performance requires continuous availability.

The rise of hybrid work, global teams, cloud collaboration, and digital communication platforms has erased many of the natural boundaries that once separated professional and personal life. Executives are accessible at all hours across multiple time zones. Strategic decisions are increasingly expected in real time. Crisis management can emerge instantly through a single cyber incident, operational failure, reputational issue, or geopolitical event.

While technology has improved organizational efficiency, it has also accelerated executive exhaustion.

The issue is not simply long working hours. Many senior leaders have always worked intensively. The deeper challenge is the absence of cognitive recovery. Continuous decision-making without adequate mental reset leads to what psychologists increasingly describe as decision fatigue, a gradual decline in the quality of judgment after prolonged periods of complex cognitive activity. For leaders operating in security-sensitive industries, this risk becomes especially significant.

An exhausted executive is more likely to miss strategic warning signs, overlook governance gaps, and make reactive rather than strategic decisions. Fatigue can gradually erode emotional regulation during crises, weaken communication consistency, and reduce the ability to maintain long-term perspective. In high-stakes environments such as cybersecurity, defense, operational resilience, and critical infrastructure, even relatively small leadership errors can create substantial downstream consequences. This is why executive well-being can no longer be viewed as separate from organizational risk management.

Leadership in the Age of Perpetual Crisis

Modern executives are not managing isolated disruptions. They are operating within overlapping crises.

A single week may involve responding to cyber threats, navigating regulatory changes, addressing AI governance concerns, handling supply chain instability, leading digital transformation initiatives, managing workforce expectations, responding to media scrutiny, and meeting shareholder performance demands. The accumulation of these pressures fundamentally changes the nature of leadership.

In previous eras, leadership cycles often allowed periods of stabilization between major disruptions. Today, many executives experience a state of perpetual responsiveness. The psychological impact is significant.

Research across leadership and organizational psychology increasingly highlights the long-term effects of sustained high-pressure environments. Burnout, cognitive overload, emotional exhaustion, sleep disruption, and reduced strategic creativity are becoming increasingly common among senior leadership teams. Ironically, the more senior the leader becomes, the more isolated these pressures often become.

Executives frequently feel obligated to project confidence and certainty even when managing significant internal stress. In highly competitive industries, vulnerability is often perceived as weakness. As a result, many leaders internalize pressure until performance, relationships, or health begin to deteriorate.

This creates a dangerous paradox: Organizations invest heavily in technology resilience, cyber resilience, operational resilience, and business continuity, while often neglecting leadership resilience. Yet executive burnout can become a major organizational vulnerability.

Redefining Work-Life Balance for Modern Leadership

The phrase “work-life balance” can sometimes feel unrealistic for senior executives. Leadership at the highest levels is inherently demanding. Critical decisions do not always respect personal schedules. Crisis response does not operate on fixed hours. Market realities often require flexibility and sustained commitment.

For this reason, many executives reject the concept of balance entirely. However, the issue is not whether leaders work hard. The issue is whether leaders can sustain high performance without compromising long-term effectiveness, health, judgment, and personal stability.

Modern executive balance is less about equal time distribution and more about sustainable energy management. This distinction is important. Sustainable leadership requires intentional recovery systems that preserve cognitive capacity, emotional stability, and strategic clarity. High-performing executives increasingly recognize several realities:

1. Recovery Improves Strategic Thinking

Periods of recovery are not productivity losses, they are strategic investments. Neuroscience research consistently demonstrates that innovation, problem-solving, creativity, and long-term thinking improve when the brain has opportunities for rest and reflection.

Executives often solve complex problems not during meetings, but during periods of mental decompression; exercise, travel, reading, time with family, or moments away from digital noise. Organizations operating in rapidly evolving sectors such as AI, cybersecurity, defense, and technology require leaders capable of adaptive thinking. Constant overload reduces that capability.



2. Boundaries Improve Leadership Quality [h3]

Effective executives increasingly establish intentional boundaries around communication, availability, and recovery. This does not mean disengagement, it means operating with discipline.

Examples include limiting unnecessary late-night communications, structuring uninterrupted strategic thinking time, protecting recovery periods after major projects or crises, delegating operational decisions effectively, reducing meeting overload, and creating clear escalation frameworks.

Executives who fail to establish boundaries often become operational bottlenecks within their organizations. Sustainable leaders empower systems and teams rather than centralizing every decision around themselves.

3. Organizational Culture Mirrors Leadership Behavior

Executive behavior shapes organizational norms. When leaders operate in permanent crisis mode, employees often feel pressured to mirror that behavior. This creates cultures characterized by chronic stress, communication fatigue, fear-based decision-making, reduced innovation, and higher employee turnover.

Conversely, leaders who model sustainable performance often create healthier organizational environments that improve retention, engagement, and resilience. In competitive industries where highly skilled talent is increasingly difficult to retain, culture itself becomes a strategic advantage.



The Security Executive's Dilemma

Leaders in the security and defense industries face unique pressures. Cybersecurity leaders, CISOs, risk executives, compliance leaders, and technology executives operate in environments where failure can have immediate, visible consequences.

A successful cyber-attack may result in financial losses, operational disruptions, reputational damage, regulatory scrutiny, legal exposure, or national security implications. This creates a professional culture heavily oriented around vigilance.

The challenge is that constant vigilance can gradually evolve into constant hypervigilance. Over time, this mindset affects decision-making, relationships, and overall well-being. Many security leaders experience difficulty disconnecting from work, persistent alertness, fear of missing emerging threats, increased stress during digital transformation initiatives, and pressure from boards and regulators.

As AI accelerates the pace and sophistication of cyber threats, this pressure is likely to intensify. Executives responsible for safeguarding organizations must therefore learn how to balance preparedness with sustainability. The future of security leadership will require not only technical expertise but also psychological resilience.

Technology: The Solution and the Problem

Technology itself plays a dual role in executive work-life balance. On one hand, AI, automation, collaboration platforms, and intelligent analytics can significantly reduce administrative burden and improve operational efficiency.

On the other hand, constant digital connectivity creates uninterrupted cognitive engagement. Executives today receive information continuously through emails, notifications, security alerts, instant messages, news updates, social media commentary, and performance dashboards. The result is often fragmented attention and constant cognitive engagement.

Attention fragmentation reduces deep strategic thinking, precisely the capability organizations most need from senior leadership. This is why forward-thinking executives are increasingly adopting intentional digital discipline. Examples include structured communication windows, notification management, dedicated strategic thinking time, reduced meeting dependency, AI-assisted prioritization, and delegation supported by automation.

Ironically, AI itself may become one of the most important tools for improving executive sustainability.

AI can help reduce repetitive cognitive workload by assisting with information summarization, draft generation, data analysis, workflow automation, scheduling optimization, reporting preparation, content creation, and risk monitoring.

However, organizations must ensure that AI enhances human decision-making rather than accelerating unhealthy productivity expectations. The goal should not be to make executives work more, instead it is to allow leaders to focus on higher-value strategic and human-centered responsibilities.

Resilient Leadership in the Age of AI

Artificial intelligence is reshaping executive leadership itself. As automation handles more operational tasks, the uniquely human dimensions of leadership become increasingly valuable. Future leaders will be differentiated not merely by technical competence but by judgment, emotional intelligence, ethical reasoning, communication, strategic foresight, adaptability, trust-building, and crisis composure. These capabilities require mental clarity and emotional stability.

Exhausted leaders rarely excel in these areas. The emergence of AI, therefore, creates an interesting paradox: The more technology automates operational processes, the more critical human resilience becomes.

This shift is especially important in sectors involving security, governance, defense, and critical infrastructure. AI may accelerate decisions. But humans remain accountable for consequences. Executives will increasingly face ethical and strategic questions involving:

- ▶ AI governance
- ▶ Autonomous systems
- ▶ Data privacy
- ▶ National security implications
- ▶ Cyber warfare
- ▶ Digital trust
- ▶ Misinformation risks
- ▶ Regulatory accountability

Managing these issues requires thoughtful leadership, not merely technical acceleration.



Building an Executive Resilience Framework

Forward-looking organizations are beginning to recognize that executive resilience should be treated systematically rather than informally.

Just as companies invest in cybersecurity frameworks, governance structures, and business continuity planning, leadership sustainability should also be approached strategically. Several practices are emerging as critical.

Prioritizing Strategic Time

Executives increasingly require protected time for long-term thinking, reflection, innovation, industry analysis, or leadership development. Without this space, leaders become trapped in operational cycles. Organizations that consume all executive capacity through meetings and reactive demands often undermine strategic agility.

Normalizing Recovery

Recovery should not be viewed as weakness. High-performance industries such as professional sports, aviation, and military operations increasingly recognize that sustained performance depends on structured recovery cycles

Corporate leadership is no different. Executives who maintain physical health, sleep quality, exercise routines, and personal relationships are often more effective during periods of crisis.

Strengthening Leadership Teams

Resilient organizations reduce overdependence on individual executives. Strong leadership distribution improves succession readiness, reduces bottlenecks, increases agility, enhances crisis response, and reduces executive burnout risk. Organizations should build leadership ecosystems rather than hero cultures.

Integrating Well-Being into Governance

Boards increasingly discuss cyber risk, ESG, operational resilience, and digital transformation. Executive sustainability may soon become another governance issue. Forward-thinking boards are beginning to ask:

- ▶ Is leadership capacity sustainable?
- ▶ Are executives overloaded?
- ▶ Is organizational culture driving burnout?
- ▶ Are crisis-response expectations realistic?
- ▶ Is succession planning sufficient?

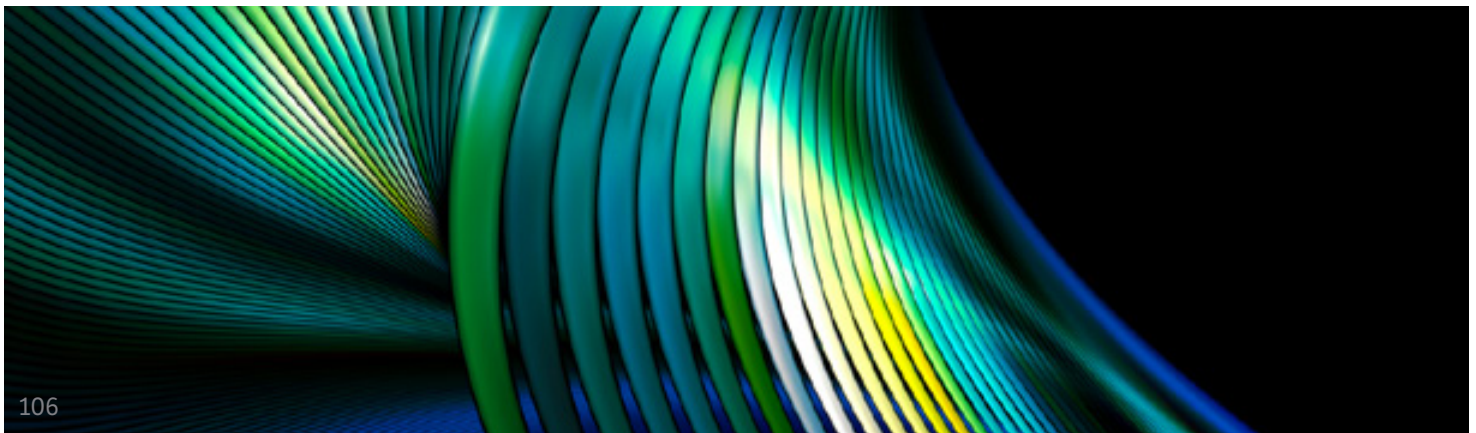
These questions are not soft management concerns. They are strategic risk management considerations.

The Human Side of Leadership

One of the greatest misconceptions in executive culture is the belief that relentless availability is synonymous with commitment. In highly competitive industries, many leaders feel pressure to remain constantly connected, responsive, and operational. Yet sustainable leadership often requires a fundamentally different approach.

Effective executives understand that leadership is not measured by perpetual activity, but by the ability to remain present, focused, and strategically composed during periods of uncertainty and pressure. The strongest leaders are rarely those who never disconnect. Rather, they are individuals capable of maintaining perspective when environments become volatile, complex, or unpredictable. In an era defined by rapid technological change, geopolitical instability, cyber threats, and constant digital communication, clarity has become a strategic advantage.

Leaders who are constantly reacting often lose the ability to think critically, anticipate long-term implications, or make balanced decisions under pressure. This perspective is especially important in today's increasingly interconnected world, where technology leaders, security executives, and governance professionals operate at the intersection of business, ethics, national security, public trust, and societal impact.





Their decisions extend far beyond organizational performance alone. They influence employees, customers, critical infrastructure, and broader communities. Such responsibility requires leaders who are psychologically sustainable. Executives who neglect personal resilience may eventually compromise professional effectiveness, even when operating with the best intentions.

Over time, chronic fatigue and cognitive overload can weaken communication, impair strategic thinking, reduce crisis management effectiveness, and erode team trust. Conversely, leaders who actively cultivate resilience often strengthen the very qualities that define effective leadership in the modern era.

They communicate more clearly, navigate crises with greater composure, foster stronger organizational trust, and maintain the strategic perspective necessary to lead through complexity. In industries shaped by constant transformation and rising uncertainty, resilience is no longer merely a personal wellness consideration. It has become a leadership capability in its own right.

The Competitive Advantage of Sustainable Leadership

The future of leadership will not be defined solely by speed. It will be defined by sustainability. Organizations capable of maintaining long-term adaptability under pressure will outperform those trapped in continuous reactive cycles. This applies equally to governments, corporations, technology firms, financial institutions, defense organizations, and critical infrastructure operators.

The leaders who shape the future will likely share several defining characteristics, including emotional resilience, strategic patience, adaptability, ethical clarity, the ability to manage uncertainty, and sustainable performance habits. These are not merely personal traits, they are competitive advantages.

In industries driven by innovation and disruption, sustainable leadership becomes a form of organizational resilience. Companies increasingly invest millions into cyber defense technologies, AI systems, and operational continuity. Yet one of the most important resilience investments may be ensuring that leaders themselves remain capable of performing effectively under sustained complexity.

Conclusion

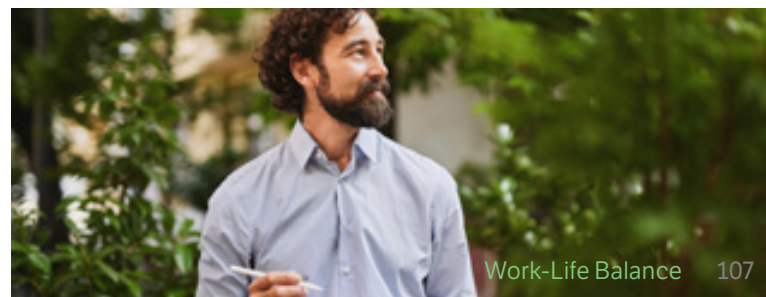
As technology continues to reshape security, governance, defense, and business itself, executives face a defining challenge: How can leaders remain agile, innovative, and resilient in environments characterized by constant disruption?

The answer will not come solely from technology, it will also come from rethinking leadership sustainability. Work-life balance is no longer a secondary conversation. It is directly connected to strategic performance, risk management, organizational resilience, and leadership effectiveness.

The future will demand leaders capable of making complex decisions in uncertain environments while maintaining ethical clarity, emotional stability, and long-term vision. Organizations that recognize this reality early will gain significant advantages; they will build healthier cultures, they will retain stronger talent, they will improve resilience.

And they will create leadership environments capable of navigating the complexity of the next decade. In a world shaped by rapid technological advancement, rising security concerns, and continuous transformation, the most valuable leadership capability may ultimately be the ability to sustain high performance without sacrificing the human foundation that makes effective leadership possible.

The future of work is not simply about smarter technology, it is about smarter, more sustainable leadership, and that may become one of the defining competitive differentiators of the modern era.



WEBINAR LIVE



Mark Your Calendar for August's Webinar

Understand how AI can enhance compliance efforts for the NIST Cybersecurity Framework, ISO/IEC 27001, and CMMC, and gain insights into the strategic opportunities and potential challenges of integrating AI into compliance programs.

**TOPIC: AI Governance and Risk Management:
Implementing ISO/IEC 42001 in Practice**

**August 27, 2026,
3:00 - 4:00 PM CEST**

Presenter:

Alexis Hirschhorn
CEO at Acuna GRC

[REGISTER HERE →](#)

Code is the New Compliance: From Checkbox to Continuous

Compliance as Code (CaC) sees organizations moving away from manual, error-prone methods.

It is the translation of regulatory policies into code, enabling the automation of compliance and the creation of a continuously monitored and enforced state. This reflects a broader movement away from manual work and toward a form of compliance that is versioned and easily updated by a compliance architect.

CaC can, and should, be integrated into the software development life cycle (SDLC). It thereby becomes a proactive, shift-left state. CaC ensures that code, infrastructure, and configurations are tested for compliance and rolled out to automate the compliance process while also providing proof of compliance.

A checkbox approach to compliance has been the approach of many organizations, and it has failed. Despite all the technology, training, and frameworks available, many organizations still treat compliance as a checkbox, despite the number of breaches and reputational damage it caused. This attitude goes by a, “We did it, we have this,” approach.

Technology is constantly struggling to catch up to protect the organization and the data subject. Compliance is a checkbox to tick rather than a competitive and strategic differentiator that should be embraced. Fines, depending on size, are seen as an expense. For this to be corrected, it starts with being driven from the top down.



What if we could take lessons from software development and apply them to compliance? The same way that software goes through a structured process to develop an application with version control, testing, and maintenance, why not use code for compliance? Apply code to operations, regulations and security.

Organizations have policies that support this which they manually enforce, reviewing and auditing systems to ensure that they are being adhered to, but this is time-consuming. What if there is a breach because a solution was not audited in time and remediation was not applied?

Operational Requirements

Operational requirements seek to create and maintain a highly available, resilient, and efficient environment. Using Infrastructure as Code (IaC) aids this by creating versioned scripts that allow for the building of environments. Instead of a manual approach to building environments, scripts can aid in ensuring that the environment is in line with the policy. Infrastructures that fail this are unstable, insecure environments that are not desirable in production.

Scripts can ensure that encryption settings are implemented as required by the organization, least privilege is implemented, logging and monitoring, and identity and access management is enforced for endpoints, ports, and default settings. This is a great use case for IaC as it would support CaC in that PCI DSS which mandates encryption can be assured. Using CaC tools, scripts can test to ensure that this is adhered to as well as enforce it by denying. These scripts can then be versioned and audited to provide audit evidence.

Using log retention as an example, IaC ensures where logs are stored, how long for, and when they should be deleted. CaC proves this and ensures that storage settings are enforced. By implementing the policy in code, it ensures that logs are stored for the organization's specified period. When combined with data tagging, PII can be protected by ensuring that those logs do not contain PII, and if they do, a level of security can be applied.

Another example is using Terraform to define the servers, then using OPA as an automated compliance check to validate against the policy defined by the organization.

Before going on, note that IaC builds the environment and CaC ensures that it meets regulations. Therefore, they are complementary. IaC executes, and CaC verifies. Think of it that way.

Security Requirements

CaC can block the accidental embedding of cryptographic keys, API tokens, passwords, and database credentials making it into the source code, especially when added to the CI/CD pipeline. As more organizations move to IaC, it can also review these and block or report the embedding of violations of the policy.

This removes the manual process of searching thousands of lines of code, which can span multiple projects, to ensure compliance with the policy. Why not use code to do this and automate the process? YAML is a widely used markup language for configuring environments. One misplaced password can open the door for threat actors.

Regulatory Requirements

SOC 2, HIPAA, AML, PCI DSS, and data privacy are among the many regulatory requirements that an organization can be subject to. Creating a regulatory crosswalk would help in the management of these. It is the foundation for a compliance architect. The compliance architect uses CaC to ensure regulatory compliance by embedding these requirements as policy and using OPA to enforce them.

A great tool for writing compliance logic is Open Policy Agent (OPA). OPA uses Rego to ensure compliance. Instead of the compliance architect defining or translating the regulations and policies and then communicating them to developers and infrastructure teams, the policy is defined using Rego and enforced with OPA.

It becomes the source of truth. These policies can be very complex or quite simple. Writing the policy in Rego means they can be versioned, moved from development to test to production, and follow the same SDLC process as any other software.

The process involves the application responding to user action, generating a JSON payload which is sent to OPA where it is evaluated, and returning a JSON response. Instead of relying on the developer's code to enforce the policy, OPA does it.

OPA's response is either approval or an error. Errors from OPA should be treated as fail so, should OPA experience an issue, it should send back fail ensuring that the policy is enforced by default. Nothing gets by. A key benefit of OPA is that it can pull policies from a server, ensuring it has the latest policy to work with. OPA is not doing the blocking or denying. Its response tells the application or process what to do.

There is then no logic for access control in the developer's application but by being controlled by CaC. This means that not only is the latest policy according to regulation, but rule is in control. Should the regulation change, OPA can run using the latest policy.

OPA can run using an older policy if there is an issue with an updated policy. The logs produced and explanation of the logic used will satisfy proof for auditors. No more screenshots need to be used as proof. No more last-minute, time-consuming rush is required to meet auditor requests.

I am not going to make this about OPA, nor any other alternative. The focus is on the benefits of CaC, and the value it brings to enforcing and evidencing compliance. This article argues that with CaC, organizations can move from a checkbox stance, a manual, reactive approach to one that is proactive, versioned, capable, and automated.

Given the constant regulatory environment that organizations operate in and the ever-increasing number of audits, implementing CaC can result in meeting the compliance needs. Some areas where CaC creates benefits include:

- ▶ Default encryption settings
- ▶ Network controls
- ▶ Identity and access management
- ▶ Cross-region replication
- ▶ Open-source license compliance
- ▶ Detection of vulnerabilities in third-party libraries
- ▶ Checking for SQL injection and cross-site scripting
- ▶ Checking for hardcoded passwords and API keys

These are just a few of many benefits which support the use of CaC in the environment. Another area of great importance is that of data privacy and the data life cycle. The latter includes data masking, data classification, and subject access requests among others.

Code can mask data by applying a policy when data is migrated from production to test environments. It can substitute customer data, such as names with masked data, or add noise to the data.

This ensures that all data has the appropriate privacy techniques applied (GDPR, among others) and provides peace of mind, especially when it comes to confirming data subjects' access rights. It can change Social Security numbers from valid numbers to made up numbers, and these changes can be dynamic or static.



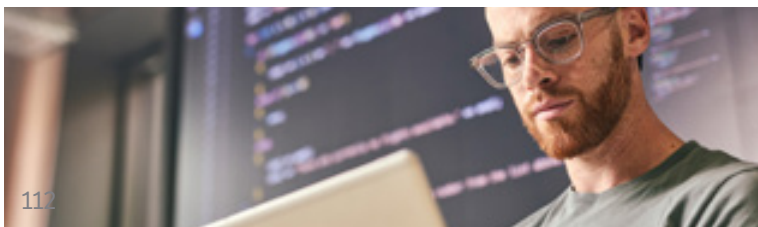
Application to the SDLC Process

Requirements Phase

In the design phase, requirements might include encryption, no public access to secure data, centralized log storage, and more. These can all be handled by CaC where policies are implemented as code that runs whenever code is promoted to another environment, such as Development to Test.

It is this phase that governance requirements are defined. Using CaC, these requirements can be captured in machine-readable policies. OPA then can be used to ensure that they are incorporated into the design before a line of code is written.

Creating a privacy architecture and then embedding that in code would ensure that no code is moved to production unless it passes the identified tests. These can then be run continuously to ensure adherence. This helps with auditing as well as the results can be pulled at any time for review.



Design Phase

Rego file for ensuring use of approved protocols:

```
package sdlc.design

# Define a list of allowed secure protocols
secure_protocols := {"HTTPS", "TLS", "SSH"}

# Deny the design by default if this rule triggers
deny[msg]
{
    # Get list of planned services in the input file
    service := input.services[_]

    # If the planned protocol is not in the list
    not secure_protocols[service.protocol]

    # so, generate a descriptive architectural violation error
    msg := sprintf("DESIGN ERROR:
    ..... Must use one of: %v.", [service.
name, service.protocol, secure_protocols])
}
```

This proves that developers cannot use unapproved protocols. This script is evidence of this and can be tested in front of the auditors by using a test input with invalid protocols.



Development Phase

Analyze code to ensure that there are no hardcoded passwords, although this should be a given. Ensure that there are no situations where there could be data leaks. Use code-testing tools to test the application while it is running.

Is encryption being used? Should role-based access control or attribute access control be implemented? What security measures have been implemented? What type of logging should be used? Is PII being stored in logs? Is there enough protection around log files such that they cannot be tampered with?

Combining Ansible with IaC, code can be scanned for compliance and violation of policies. These include overly permissive identity and access management roles. Having developers run this prior to migrating code to servers ensures that violations never reach the server.

Using CaC automates these tests. The output can be used to ensure adherence. CaC can also see the development of rules that ensure that no developer uses open-source software with a license that violates the organization's rules on allowed software. The output then shows adherence, since its results can be used in metrics to show the number of breaches of the policy.

Rego file for ensuring the blocking of forbidden licenses:

```
package sdlc.development

# Define list of forbidden copyleft licenses
blacklisted_licenses := {"GPL-3.0", "AGPL-3.0", etc.}

# Denial message if a dependency violates the policy
deny[msg] {
    # Loop through the dependency licenses map in the
package file
    license_type := input.dependency_licenses[package_
name]

    # Check if the license type is in the forbidden blacklist
blacklisted_licenses[license_type]

    # Failure message for the developer's terminal
    msg := sprintf("COMPLIANCE FAILURE: Package .....
uses a forbidden copyleft license (%v). Proprietary apps cannot
use GPL/AGPL dependencies.", [package_name, license_type])
}
```

For Privacy

```
Package authz
Default allow:=false
```

Deny by default

Testing Phase

Items developed should be tested here, including full functionality, privacy, and security. User testing feedback should be fed back into development to ensure that their requirements and issues have been addressed. Vulnerability scanning should also be performed. With a shift-left mentality, the aim is to implement privacy protection early in the phases as it gets more expensive to “bolt on” later. CaC, with policies written as rules that are machine-readable, ensures that testing confirms that what will be shipped conforms to the rules/policies of the organization.

Rego file for issues pertaining to vulnerabilities flagged as high or critical. This works with automated testing. Depending on the results, high-ranking vulnerabilities will be flagged with a message and again, this is strong audit evidence.

```
package sdlc.testing

# 1. Define the severe times that should block a release
blocking_severities := {"HIGH", "CRITICAL"}

# 2. Denial message if any test failure matches the criteria
deny[msg] {
    # Extract each individual vulnerability from the test
    results array
    vuln := input.vulnerabilities[_]

    # Check if the vulnerability severity is listed in our
    blocking set
    blocking_severities[vuln.severity]

    # Generate an explicit compliance failure message for the
    test logs
    msg := sprintf("QUALITY GATE FAILURE: Testing
    found a %v vulnerability (%v'). Fix this flaw before merging
    code.", [vuln.severity, vuln.name])
}
```

What makes Rego and OPA valuable is that they not only produce reliable, provable evidence but are also well suited for testing. See this example taken from [Open Policy Agent](https://www.openpolicyagent.org/docs/policy-testing):

```
package authz

allow if {
    input.path == ["users"]
    input.method == "POST"
}

allow if {
    input.path == ["users", input.user_id]
    input.method == "GET"
}

(https://www.openpolicyagent.org/docs/policy-testing)
```

Testing it:

```
package authz_test

import data.authz

test_post_allowed if {
    authz.allow with input as {"path": ["users"], "method":
    "POST"}
}

test_get_anonymous_denied if {
    not authz.allow with input as {"path": ["users"], "method":
    "GET"}
}

test_get_user_allowed if {
    authz.allow with input as {"path": ["users", "bob"],
    "method": "GET", "user_id": "bob"}
}

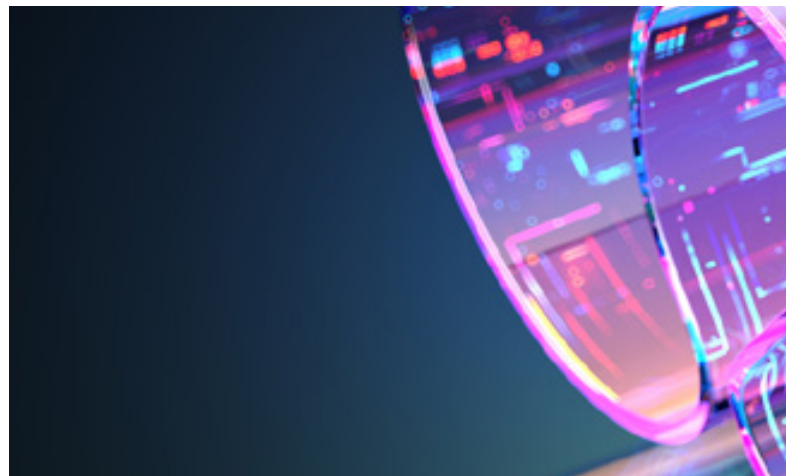
test_get_another_user_denied if {
    not authz.allow with input as {"path": ["users", "bob"],
    "method": "GET", "user_id": "alice"}
}

(https://www.openpolicyagent.org/docs/policy-testing)
```

Auditors often ask for proof, particularly when an organization states, “We only allow...” or “If the system is given, it would...” Create a test file with the values auditors suggest and show them the results. Moving authorization to a policy tool like OPA saves hundreds of lines of code, along with significant coding and regression testing.

Deployment/Maintenance Phase

As applications change along with the business, it is a prime location for policies to be violated. CaC ensures that through the automation of policies, this does not happen. CaC can also ensure that the built environment conforms to the prescribed policy.



Continuous monitoring, supported by automation, can provide evidence of conformity. Logs from the automation can be reviewed by internal and external auditors.

Rego file to block the deployment of code if it is to be done during a critical time for the business:

Logic:

- ▶ Define the no-go dates.
- ▶ Determine whether the target is a production system.
- ▶ Check the date of deployment against the dates.
- ▶ Check if an emergency fix that must go in.

```
package sdlc.deployment
```

```
# 1. Define a set of blacklisted corporate freeze dates
```

```
freeze_dates := {  
    "2026-12-26", # Boxing Day  
    "2026-12-31", # New Year's  
    "2026-12-25" # Christmas Day  
}
```

```
# 2. Denial if the deployment violates operational rules
```

```
deny[msg] {  
    # Check if the target environment is production  
    input.environment == "production" #we want this  
    only for production
```

```
    # Check if the current deployment date falls within a  
    blackout dates  
    freeze_dates[input.deployment_date]
```

```
    # Allow an exception ONLY if it is explicitly flagged as an  
    emergency hotfix  
    input.is_emergency_fix == false
```

```
    # Generate the operational compliance block message  
    msg := sprintf("DEPLOYMENT TO PROD BLOCKED:  
Target date '%v' is during a production blackout period. Routine  
changes are prohibited. Emergency fixes are allowed. Ticket  
'%v' aborted.", [input.deployment_date, input.change_ticket])  
}
```



Providing Proof and Metrics

Each time an application queries OPA, it can generate a JSON log. This log serves as evidence, showing the policy, what was inputted, the result (proof that the policy is being used), and the path to the relevant Rego file, package, and the rule that was executed. OPA also produces output to prove to the auditors that it was running, and it can show that high traffic for instance did not affect its performance.

In short, compliance does not have to be a checkbox affair. With tools like Terraform and OPA, it can be automated, provable, and up to date with most, if not all, compliance requirements. With CaC, organizations have a way of converting their policies into machine-readable code. Then implement, version, and use that code to ensure that most of their requirements are met.

CaC, combined with complementary IaC, sees compliance move to a proactive, audit-friendly state where being in compliance and proving it is not only ideal but also the standard to strive for.



```

def insert_record(self, key, value):
    timestamp = datetime.now().strftime("%Y-%m-%d %H:%M:%S")
    cursor.execute(
        """INSERT INTO records (key, value, timestamp)
        VALUES (?, ?, ?)""", (key, value, timestamp))
    connection.commit()

def fetch_records(self):
    cursor.execute("""SELECT * FROM records
    ORDER BY timestamp""")
    return cursor.fetchall()

// Common method (template method pattern)
void runTransaction(const std::string& query) {
    if (connect(default_connection)) {
        std::cout << "[Transaction] Starting...\n";
        executeUpdate(query);
        disconnect();
        std::cout << "[Transaction] Completed.\n";
    }
}

virtual std::string getType() const = 0;
};

class PostgreSQL : public Database {
public:
    bool connect(const std::string& connection_string) const {
        // ...
    }
};

```



Nigel Pierre

Compliance Architect at Sagikor

Nigel Pierre is a compliance architect, sitting between technology and compliance. Acting as a bridge, Pierre ensures the organization can meet its regulatory requirements. He is well-versed in data privacy across multiple jurisdictions, AML, FATCA/CRS, technology compliance, and controls. He is also a software engineer with a unique perspective on the role of technology's impact on cybersecurity, data privacy, risk management, auditing, and controls. Pierre has worked with internal and external auditors and regulators, ensuring the translation of regulatory requirements into controls and solutions. As his organization is multi-jurisdictional, he has to learn and understand as well as appreciate the nuances of the appropriate regulations.



THE EXPERT

Silence by Design: Why Good People Stay Quiet About Breaches

Fifty-eight percent of security professionals have been told to keep a breach confidential when it was required to be reported. Every reporting deadline in Europe starts at a precise moment, and it is not the one you would expect. It is not in your logs. It is in a conversation that may never happen. So what would it cost somebody in your organization to tell you something you would rather not hear?

Have you ever been asked to keep an incident quiet? I am not talking about a cover-up, nothing as dramatic as that. Just handled internally, kept in the room, managed discreetly while somebody more senior decides how serious it really is. Have you ever watched an incident being logged as something smaller than it actually was, so that it would draw less attention than it deserved? Or sat in a meeting where somebody wondered aloud whether an event really met the reporting criteria, in a tone that made clear which answer was wanted? If you have, you are in the majority.

[Bitdefender's 2025 Cybersecurity Assessment Report](#), based on 1,200 IT and security professionals, found that 58% had been told to keep a breach confidential even when it was required to be reported. Among C-level respondents, the figure was 69%. In the United Kingdom it was 58%, in Italy 53%, in Germany 48%. The figure is up 38% on 2023.

That was my first question. Here is the second. Without looking it up, just off the top of your head, how long do you have to report a significant incident under NIS2? Not roughly. Exactly. And from what moment does that period begin to run?

And my third question, and the one that should worry you most. If somebody on your team made a serious mistake with customer data this morning, would you know about it before lunch? Would they tell you at all? And how long could they keep that decision to themselves? A week? A month? Until the next audit?

One company already knows its answer, and it was longer than you would think.

A Year of Silence

On 4 November, 2016, Uber's chief security officer testified under oath to a regulator about how the company protected its data. Ten days later, on 14 November, criminals emailed him to say they had taken the personal data of 57 million Uber users, including around 600,000 driver's license numbers. Both dates are on the record in the [US Department of Justice account of the case](#).

The chief security officer did not report it. Uber paid the criminals 100,000 dollars in bitcoin and had them sign non-disclosure agreements stating that no data had been taken. Yes, you read that right. And the breach became public a year later, on [21 November, 2017](#), when Uber's new chief executive published a statement acknowledging "our failure to notify affected individuals or regulators last year."

Then the regulators arrived. The Dutch data protection authority [fined Uber 600,000 euros](#) because "it did not report the data breach to the Dutch DPA and the data subjects within 72 hours after the discovery of the breach." The UK Information Commissioner's Office [fined it 385,000 pounds](#) over a breach affecting 2.7 million UK customers and 82,000 UK drivers.

The prosecutors came next. In October, 2022, a jury convicted the chief security officer of two federal crimes: obstructing a government investigation and hiding a crime he knew had been committed. The DOJ recorded what he had told a subordinate at the time. They "can't let this get out." The information needed to be "tightly controlled." And outside the security group, the story was to be that "this investigation does not exist."

If that sounds like old news, look at the dates. He appealed and lost: in March, 2025, a US federal appeals court [upheld his conviction](#). He then asked the US Supreme Court to hear the case, and in [June, 2026](#), it refused. The conviction is final. There is nothing left to appeal.

And before Uber, Joseph Sullivan spent [eight years at the US Department of Justice](#). He was a prosecutor in the US Attorney's Office for the Northern District of California and a founding member of its unit dedicated to technology crime. In [December, 2001](#), that office named him as one of the Assistant US Attorneys prosecuting the first case ever brought under the Digital Millennium Copyright Act. More than twenty years later, the same office convicted him. Nobody understood these laws better than he did, and he broke them anyway. So if knowing the rules is not what makes somebody speak up, what is going to make the person on your team speak up tonight?

Now look at what actually failed. No sensor was missing, no control was absent, and nobody needed a better tool, a bigger budget, or a more mature framework. One person knew within hours, and the organization took a year to say so. Every euro and every pound of those fines was charged on the gap between knowing and telling.

So when should Uber have reported the breach? When did the clock start ticking exactly?

Your Deadline Does Not Start When You Think It Does

Here is the answer to my second question. Check it against whatever you had in mind.

[NIS2 Article 23\(4\)](#) requires an early warning "without undue delay and in any event within 24 hours of becoming aware of the significant incident," and then a fuller notification within 72 hours. [GDPR Article 33\(1\)](#) requires notification

“without undue delay and, where feasible, not later than 72 hours after having become aware of it.”

Look at the word “and” in both quotes. Each law asks for two things, not one. First, report as soon as you reasonably can. Second, never take longer than 24 or 72 hours, whatever your reason. Most of us know the numbers. Fewer know when they start counting. And almost nobody thinks about the first rule, which is the one that catches people. You can report at hour 70 and still be in breach, if you could have reported at hour 5.

Both clocks begin when your organization “becomes aware.” Not when the attack happens. Not when the log entry is written. Not when the SIEM correlates the alert. When your organization becomes *aware*.

The European Data Protection Board has been careful about what that means. In [Guidelines 9/2022](#), a controller becomes aware when it “has a reasonable degree of certainty that a security incident has occurred that has led to personal data being compromised.” You are allowed to have a short period of investigation first. But the same guidelines add a requirement that most organizations skip past: “the controller should, therefore, have internal processes in place to be able to detect and address a breach.”

In a [December, 2020, decision against Twitter](#), the Irish Data Protection Commission put it more directly. Awareness “must be understood in the context of the broader obligation on a controller to ensure that it has appropriate measures in place to facilitate such ‘awareness’”

Read that twice, because it is the whole argument. A regulator has told you that *your awareness* is your responsibility to construct. You cannot claim the clock never started because nobody told you, if the reason nobody told you is that you never built anything that would make them want to, or have to.

So your awareness is your responsibility, and your escalation procedure is written down, reviewed, and signed off. Both of those things are true at the same time. Neither of them tells you what happens at six in the evening on the Friday before a public holiday, when the person who would have escalated has already gone home. That is where the gap opens.

The Actual Gap in Your Breach Disclosure Process

Every breach disclosure process has a gap in it. The document is usually fine. The problem is the distance between what it says and what people do when nobody is watching.

The Irish Data Protection Commission measured one of those gaps and published it. It [fined Twitter 450,000 euros](#) for reporting a breach late and for failing to document it properly, and the [decision itself](#) records every step and the date each one happened.

A bug bounty report arrived on 26 December, 2018. A ticket was raised three days later. Then nothing, because “as a result of the winter holiday schedule the internal Twitter security team did not review the 29 December ticket... until 2 January, 2019.” Legal spotted the GDPR problem on 3 January. An incident ticket was opened on 4 January, and the Data Protection Officer was not added to it. The DPO finally heard about the breach on 7 January, “(verbally) of the Breach during a Twitter Group weekly team meeting.”

Twelve days. That is how long it took for the news to reach the one person legally responsible for telling the regulator. He heard it in a weekly meeting, out loud.

The regulator’s explanation was one sentence. “5/Escalation to Legal’ was not followed as prescribed in the runbook, and resulted in a delay in notifying the Global DPO.” That “5/” is step five of Twitter’s own written procedure.

Every step in that chain was taken by somebody doing their job as they understood it. Nevertheless, a holiday schedule and one skipped step cost 450,000 euros.

Then Capita. In October, 2025, the UK’s data protection regulator (the Information Commissioner’s Office) [fined the outsourcing group 14 million pounds](#) over an incident affecting 6.6 million people, finding that “a high priority security alert was raised within ten minutes of the breach, but Capita took 58





The Friday Evening Problem

Let us put my first and third questions back together. Have you ever been asked to keep an incident quiet? And would somebody on your team tell you if they made a serious mistake with customer data?

One is silence handed down to you. The other is silence kept from you. Between them sits a position we have engineered for the security profession, and it is genuinely unreasonable. You are more likely than not to be asked at some point to stay quiet. And you have never been more personally exposed if you do.

[Splunk's CISO Report 2026](#), based on 650 CISOs across nine countries, found 78% concerned about being held personally liable for security incidents, up from 56% a year earlier. Under NIS2, Article 20(1) holds management bodies liable for infringements, and Article 32(5)(b) allows a competent authority to seek a temporary ban on an individual exercising managerial functions in an essential entity.

To be fair, far more people fear this than have ever been punished for it. I went looking for a European manager who has actually been banned or personally fined under NIS2. There is not one in the published record. And supervisory authorities do not have to publish their fines, so that is not proof that none exist. But there is a blunter reason the record is thin. On [8 July, 2026](#) the European Commission referred Ireland, Spain, France and the Netherlands to the Court of Justice for failing to transpose NIS2 into national law at all, more than twenty months after the deadline. The power exists on paper. In four member states, not even that. And none of it is comforting when you are the person holding the information on a Friday evening.

hours to respond appropriately, against a target response time of one hour.”

Ten minutes to detect. Fifty-eight hours to act. In between lay more than two days in which Capita's own systems had already flagged the problem and nobody had dealt with it.

That is the gap, and almost nothing measures it. We audit whether the monitoring exists. We audit whether the procedure exists. In many years of auditing management systems I have never once had to test a control that measures how long the alarm sat there before a person moved, because in the frameworks we certify against, there is not one.

Notice what Twitter and Capita have in common. Nobody was hiding anything. No career was at risk, no awkward conversation was being avoided, and nothing was at stake for any of the people in that chain. It broke anyway, on a holiday schedule and an unread alert.

Now picture the same chain on a Friday evening, with one person in it who knows that speaking up is going to cost them something.



Now consider the person several levels below you, the one who clicked the thing, or used the tool they were told not to use, or copied the file somewhere convenient. They face the same calculation with none of your standing and none of your context. Research from Keeper Security found that 41% of cyber-attacks were never disclosed to internal leadership, and the reason given most often was [fear of repercussion](#), at 43%.

And the law does nothing for them. NIS2 Article 23(1) provides that “the mere act of notification shall not subject the notifying entity to increased liability.” GDPR Article 83(2) requires supervisory authorities to weigh “the manner in which the infringement became known.” Both protect the organization in its dealings with the regulator. Neither protects an employee from their own employer, and neither stops an internal report from being used to decide who gets blamed.

So everyone in this chain is behaving rationally. The executive who keeps it contained and the analyst who says nothing are both doing the sensible thing, given what it would cost them not to. That is what makes this a design problem rather than a discipline problem. And that is the good news, because a design is something you can change.



What This Looks Like When It Works

Let me be clear about what I am not asking for. I am not asking you to go soft on genuine misconduct, or to forgive everything. The organizations that get this right are not more lenient than yours. They are better informed, and they are better informed because their people have stopped calculating.

You would notice the difference within a quarter, and not in your dashboards. You would notice it in how you find things out. Somebody comes to you on a Tuesday afternoon with something small and awkward, and they do it without rehearsing the sentence first. The alert that fires at ten minutes is answered in twenty, because the person who saw it was not wondering whose problem it was. Your Data Protection Officer learns about a breach from a ticket with their name on it, on the day it opens, rather than out loud in a weekly meeting twelve days later.

And when the notification does go to a regulator, it reads differently. It becomes a considered account of what happened and what you did about it, written by people who had time, instead of an explanation of why it took so long. Your incident log fills with things your people told you rather than things you eventually discovered. Your auditors stop asking careful questions about the gaps in it.

None of that requires anybody to be braver than they are today. It requires that telling you stops being the expensive option. Four changes will do most of the work.



Start This Week

All four of the suggestions below aim at a single moment: somebody in your organization *is deciding* whether to tell you something awkward. Each one makes it more likely that they will. Get that moment right and the deadlines look after themselves.

Get it wrong and no framework will help you, because every framework you certify against quietly assumes that somebody spoke up.

- 1. Write down where honest error ends and misconduct begins.** Two sentences pinned to your intranet do more than a policy nobody opens: “Clicking a phishing link is not a disciplinary matter. Not telling us you clicked one is.” And be clear about which side of the line concealment sits on. The distinction that matters is not between large mistakes and small ones. It is between making a mistake and hiding one. An honest error made while doing the job is a process problem. Deciding not to mention it is not. Write down the consequences attached to each.
- 2. Make reporting quicker than staying quiet.** Capita gives you the benchmark and the warning in one line: a one-hour target response time, and 58 hours in practice. Do the same arithmetic on your own reporting route this week. Count the clicks from “I think I made a mistake” to “somebody has acknowledged it.” How many clicks does it take to say nothing? Zero. That is what you are competing with. Get it down to one channel, one step, and an acknowledgement the same day.
- 3. Assume something has not reached you yet.** Zero trust architecture begins by assuming an attacker is already inside rather than waiting for proof, and disclosure deserves the same posture. So stop asking whether anything is being withheld and go looking for it. Take your last ten incidents and write down how each one actually reached you. A colleague walked in. A tool fired. A customer complained. An outside researcher emailed. If most of them arrived from outside your organization, or through a corridor conversation rather than your official channel, then your reporting route is not what is finding your incidents, and whatever it is missing is still out there.
- 4. Count how many incidents a person reported rather than a tool detected.** At Uber a person knew within hours, and the tools never mattered. At Capita a tool knew within ten minutes, and no person acted for 58 hours. Tag last year’s incidents “human” or “tool,” and then look at how long each type took to reach somebody who could act. If your human column is nearly empty, resist reading that as proof your tools are excellent. It more likely means nobody is telling you anything.

Then hold one meeting. Invite security, HR, legal, and whoever owns incident response. Put this on the agenda. *If somebody here made a serious mistake with customer data this morning, what exactly would happen to them for telling us? And would they have known that answer before they decided whether to tell us?*

If the room cannot answer both questions quickly, and give the same answer, you have just found the reason your reporting deadlines are theoretical.

We secured the systems. We secured the process. We secured behavior, and then we started securing the AI agents. We never secured the ten seconds in which somebody decides whether to tell us about something we would rather not know. That is where your 24 hours begins. That is where your 72 hours begins. Not in the SIEM. Not in the runbook. In one person’s head, at the end of a long day, weighing what honesty is going to cost them.

Silence by design is still a design. Change it.







Nelson Marques

Cybersecurity Officer at VodafoneZiggo

Nelson Marques is cybersecurity officer at VodafoneZiggo and a PECB Certified Trainer, delivering ISO/IEC 27001 Lead Auditor and Lead Implementer courses alongside other industry-leading certifications such as CISM and CCSP.

Across more than two decades in hospitality, telecoms, financial services, and higher education, he has held senior security roles at organisations including Marriott, Kyocera EMEA, N26, and Vodafone, leading global governance, risk, and compliance functions across European and North American operations.

He delivers training courses and online webinars on cybersecurity matters in English, Portuguese, and Spanish, including AI Risk Management covering the NIST AI Risk Management Framework, ISO 31000, the EU AI Act, and the OWASP Top 10 for AI.

He speaks regularly on cybersecurity and AI, and has always strived to translate complex security challenges into plain language that boards, engineers, and everyday employees can act on.

His work is shaped by two ideas: that “security is baked in, not sprayed over,” and a discipline he calls “Culture by Design,” which holds that security culture deserves the same intention, rigor, and engineering as system architecture itself.



PECB CONNECT

Connecting Organizations with Trusted Auditors

Finding the right auditor shouldn't be a challenge.

PECB Connect is a global platform that brings organizations and qualified auditors together, making it easy to connect based on expertise, qualifications, and industry-specific needs.

Organizations can quickly identify verified professionals for certification, compliance, or internal audits, while auditors can showcase their expertise, access global audit opportunities, and grow their professional network. With streamlined collaboration, advanced search capabilities, and a centralized platform for managing auditor relationships, PECB Connect helps simplify the auditing process from start to finish.

Discover how PECB Connect can support your organization's auditing needs or advance your auditing career.

[READ MORE →](#)

Risk-Informed Decision Making: Helping Leaders Navigate Digital Uncertainty

Modern leadership moves at a pace that often leaves
reflection in the dust.

Modern executive teams are under continuous pressure to innovate, deploy, and scale digital capabilities faster than the competition. Whether it is adopting enterprise AI models, migrating core infrastructure to cloud architectures, or automating legacy workflows, the mandate is clear: move quickly or risk obsolescence.

However, this pursuit of speed can create an uncomfortable paradox. The very digital initiatives designed to create enterprise agility often end up introducing unchecked complexity, invisible dependencies, and cascading operational fragilities. Leaders find themselves driving at top speed down a foggy highway, celebrating their momentum while quietly praying that nothing appears unexpectedly in their lane.

This article is not a lecture on pump-the-brakes conservatism. Risk professionals who exist merely to say “no” are quickly bypassed or relegated to rubber-stamp compliance functions. Rather, this is a practical exploration of how executive decision-makers can transition from treating risk as a governance hurdle to using risk insights as a steering mechanism; enabling calculated and confident velocity in an unpredictable landscape.

Understanding the Exposure: Digital Friction and Executive Blind Spots

Before leadership can make risk-informed decisions, we must be brutally honest about why traditional decision-making models break down during digital transformation. The failure is rarely due to a lack of talent or intelligence; it stems from structural mismatch between how decisions are framed and how risk actually manifests in complex systems.

Digital uncertainty carries a specific, compounding set of executive blind spots:

► **Coupled Complexity**

Strategic choices are rarely isolated anymore. A seemingly minor vendor integration or cloud API configuration can ripple across cybersecurity, regulatory compliance, data privacy, and operational continuity.

Executives are often asked to approve business cases that evaluate commercial returns in depth while treating technical and operational dependencies as mere footnotes.

► **The Illusion of Historical Precedent**

Traditional risk assessments rely on historical data to project future likelihood but when deploying novel architectures or AI frameworks, historical precedent does not exist.

Relying strictly on lagging indicators gives boardrooms a false sense of security while forward-looking, tail-risk exposures compound unnoticed.

► **Information Asymmetry and Optimism Bias**

When project teams are incentivized purely on delivery timelines and go-live milestones, risk reporting naturally suffers from filter drag.

Red flags are diluted as they travel up the management chain, transforming from critical system vulnerabilities into mild, amber-flagged project risks by the time they reach executive dashboards.

► **Accountability Without Visibility**

Senior leaders face immense accountability for operational failures, yet they are often separated from the granular operational reality of their systems.

This gap breeds either over-reliance on technical experts who lack commercial context, or top-down mandates that force technical teams into unsafe shortcuts.

Recognizing the Signals: The Digital Uncertainty Register

In governance work, the most dangerous exposures are almost never the catastrophic black swan events that strike without warning. They are the subtle, normalized deviations in the form of warning flags that leadership gradually gets used to seeing and rationalizing until failure becomes inevitable.

When evaluating digital transformation initiatives, executives must train themselves to look beyond standard status reports and track the underlying operational indicators:

Uncertainty Signal	What It Looks Like at the Leadership Level
Dashboard Wash	Projects showing perpetual green statuses despite ongoing scope changes, delayed testing, or unresolved audit items.
Architecture Drift	Target operating models being bypassed via shadow IT or workarounds to hit arbitrary delivery deadlines.
Concentration Risk	Critical infrastructure dependent on single key vendors or niche internal personnel without succession redundancy.
Governance Lag	Policy frameworks and risk thresholds being updated months after new technology stack deployments.
Metric Distortion	Focusing on adoption velocity and user numbers while ignoring exception rates, error logs, and data quality flaws.
Feedback Friction	Technical teams expressing hesitation or resistance during risk reviews, signaled by compliance-driven answers.



The Decision Architecture: Strategic, Tactical, and Operational Alignment

To make risk-informed decisions, executives must avoid treating risk management as a single uniform activity. Just as in operational resilience, decision-making architecture must operate across three distinct operational layers:

- 1. Strategic Layer - Risk Appetite and Capital Allocation.** This is about asking whether an initiative aligns with the organization’s overarching risk capacity. It establishes non-negotiable guardrails—defining how much volatility, regulatory exposure, or operational disruption the enterprise is genuinely willing to absorb in exchange for growth.
- 2. Tactical Layer - Portfolio Governance and Trade-off Analysis.** At this layer, leadership evaluates competing priorities. When delivery schedules collide with risk remediation, tactical governance provides the framework to evaluate trade-offs objectively rather than defaulting to schedule pressure.
- 3. Operational Layer - Control Effectiveness and Real-Time Feedback.** This is where execution lives. It requires robust, continuous testing of controls, real-time threat monitoring, and immediate escalation channels when operational metrics exceed predefined tolerances.

When these three layers are aligned, risk management stops feeling like an administrative bottleneck. It becomes an integrated radar system that allows leadership to steer dynamically around emerging hazards.





Control Design for Executive Decision-Making

Generic advice often tells leaders to “foster a risk-aware culture” or “improve communication.” While well-intentioned, these directives lack actionable specificity. To build a truly risk-informed decision-making environment, leadership must implement concrete controls:

Risk-Adjusted Business Cases - Integrate formal scenario analysis and risk-reward thresholds directly into investment committee charters. Business cases should not receive capital allocation approval based on financial ROI alone; they must present a clear evaluation of residual risk exposures, dependency maps, and exit strategies.

Pre-Mortem Analysis - Before committing major capital or launching enterprise scale-outs, conduct structured pre-mortem sessions with cross-functional teams by asking: “Assuming this deployment fails catastrophically 18 months from now, what caused it?” This exercise breaks optimism bias and surfaces structural risks that traditional reporting misses.

Dynamic Risk appetite Triggers - Establish dynamic risk appetite triggers rather than static yearly policies. Define specific threshold metrics (e.g., error rates, customer friction points, or security vulnerabilities) that automatically trigger executive review and decision-making before an issue scales into a crisis.

Automated Key Risk Indicators (KRIs) - Implement real-time risk dashboards that pull directly from operational systems rather than relying purely on manual status decks. Automated telemetry provides leadership with objective data on system health and compliance posture.

The Executive Identity Trap: Certainty vs. Agility

There is a psychological hurdle that many executive leaders face when navigating digital transformation: the trap of demanding absolute certainty in an inherently uncertain environment.

Traditional corporate leadership models reward decisive confidence and definitive answers. In stable environments, this works well but in rapidly evolving digital contexts, expecting complete clarity before making decisions leads to paralyzed governance, delayed execution, or performative certainty where teams fabricate confidence to satisfy executive expectations.

Risk-informed leadership requires a fundamental mindset shift. It requires replacing the illusion of absolute control with structural agility and intellectual honesty. A competent leader understands that uncertainty is not a failure of planning; it is the natural operating environment.

The true test of executive capability is not whether you can eliminate all risk, but whether you can build an operating model that remains resilient and adaptable when assumptions turn out to be wrong.

Organizational Hazards: When Governance Becomes Performative

It is vital to distinguish between genuine risk-informed decision-making and performative risk management. Many organizations invest heavily in risk frameworks, enterprise software, and governance committees, yet remain dangerously exposed.

The signs of performative governance are unmistakable: risk registers that exist as static spreadsheets updated once a quarter for audit purposes; risk committees that spend hours discussing formatting and escalation mechanics rather than strategic trade-offs; and leadership cultures that celebrate compliance sign-offs while ignoring operational realities.

When governance becomes performative, it creates a dangerous false sense of security. Executive teams believe they are protected because the paperwork is complete, right up until an operational crisis exposes the underlying fragility. Recognizing and dismantling performative governance is a primary responsibility of sound leadership.

The Practitioner's Responsibility

Ultimately, risk-informed decision-making is not about slowing down innovation or avoiding bold moves. It is about enabling sustainable velocity. It is the realization that the fastest cars on a racetrack are equipped with the best brakes—not so the drivers can stop, but so they can take sharp corners with complete confidence.

Leaders who master this discipline do not treat risk management as a regulatory burden or a back-office utility. They embed risk insights into the very core of their strategic choices. By recognizing signals early, building structured decision controls, and maintaining intellectual honesty in the face of uncertainty, executive teams can navigate digital transformation with clarity, resilience, and gain genuine competitive advantage.



Hershin Marcello

CRISC

Hershin is a risk management professional with an international career spanning the financial services industry across banking, insurance, and management consulting.

Over the course of his career, Hershin has built deep expertise across operational risk, market conduct, compliance, investigations, safety management, assurance, operational resilience, and governance transformation.

He lives by a philosophy of accepted residual risk and maintains rigorous control discipline in all areas of life - with one well-documented exception involving cheesecake, which he has formally classified as a black swan event.

“If you never take any risks in life, your rewards will always be in your dreams.” - Hershin Marcello (CRISC)

If you would like to connect with Hershin on all things Risk Management, you can find him on LinkedIn by clicking the icon above.

ENRICH YOUR PROFESSIONAL

Advance further with PECB's new

Contact us at support@pecb.com

Updated Training Courses	Language	
ISO 14001 Lead Auditor	English	>
ISO 14001 Lead Implementer	English	>
GDPR - Certified Data Protection Officer	English	>
Certified Artificial Intelligence Professional	English	>
Lead Crisis Manager	English	>
ISO 22000 Foundation	English	>
Certified Incident Responder	English	>

COMPETENCIES AND CAREER

view and updated training courses!
or visit our [website](#) for more.

New Training Courses	Language	
ISO 14001 Transition	English	>
NIST Cybersecurity Foundation	English	>
ISO 37001 Foundation	English	>
Certified Cloud Security Analyst	English	>
ISO 56001 Lead Implementer	English	>

Upcoming Training Courses	Language	
Certified Artificial Intelligence Manager	English	>
Certified AI Security Professional	English	>
Certified EU AI Governance Professional	English	>

TITANIUM



PLATINUM

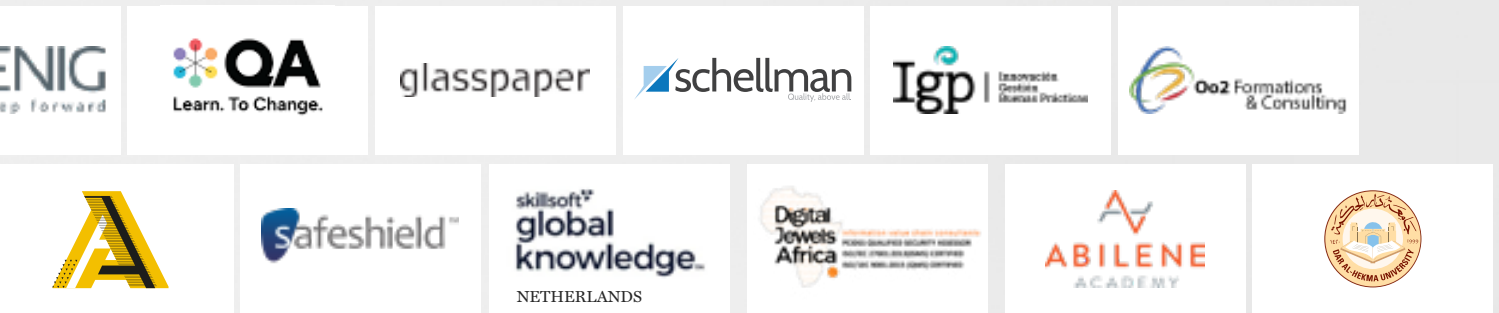


GOLD PARTNER



Note that PECB Partners are listed as per the credentials

PARTNERS



PARTNERS



PARTNERS



acquired from January 1, 2025, to December 31, 2025.



Lead with Confidence

with
PECB's Training Courses.