

Enterprise-grade Agentic AI: Secure, governed, and sovereign by design

Atos gives enterprises the confidence to scale
agentic AI safely, securely, and with full control over
cost, value, and sovereignty.



Atos

C-Suite brief

Autonomy is racing ahead. Answers aren't.

Agentic AI is moving at breakneck speed but it has left C-level leaders holding the hardest problems: where is the defensible value, how do we secure write-enabled agents, where do we draw sovereignty lines, and how do we scale without losing control or predictability? Too many pilots impress in demos yet stall at production because the industrial-grade answers (governance, security, architecture, and cost discipline) arrive last, if at all.

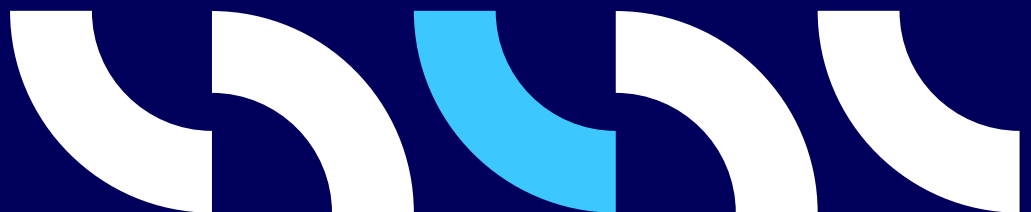
If you are a global organization, this gets real fast. Your data lives across jurisdictions with conflicting residency and access rules. Agents need write access into ERP, CRM, ITSM, and OT- expanding the blast radius if prompts, memory, or tools are manipulated. Third-party models and tools introduce supply-chain risk and low-/no-code spawns agentic sprawl duplicating logic and obscuring ownership. Boards demand auditability and accountability; CISOs need runtime guardrails, revocation, and kill switches; CFOs want unit economics and throttles as token, tool, and loop costs scale. And everything must remain sovereign where it matters, without locking you into a single vendor stack.

This whitepaper is your field guide to move from promise to production without surrendering control. It reframes agentic AI as production software and shows how to scale autonomy only when value, safety, and compliance are demonstrable.

In this whitepaper, we discuss:

- A disciplined blueprint for “Services as Software”, with Atos’s Sovereign Agentic Studio that turns workflows into software-delivered outcomes, progressively and safely
- The control plane for trust at scale: runtime governance, zero-trust access, behavioral security, kill switches, observability and immutable audit trails
- Sovereignty by design: where and how to assert control (data, models, orchestration, runtime), including air-gapped, private, hybrid and jurisdiction-aware deployments
- AgentOps & FinOps from day one: unit economics, spend throttling, and value tracking tied to business KPIs to prevent sprawl and cost surprises
- Data readiness as a precondition: curated, governed and trusted data to avoid autonomy amplifying bad inputs
- Evidence from production that shows autonomy can scale with control: reductions in MTTR and ticket volumes, more incidents resolved autonomously and lower run costs, achieved with confidence-based controls

If you are accountable for outcomes and risk, this is a pragmatic path to scale autonomy on your terms, with value, security, sovereignty and cost under control.



THE PROMISE

Agentic AI is to the modern enterprise what electricity and the internet were to previous generations:

a radical leap that redefines how organizations operate and create value. Unlike traditional automation, agentic AI consists of systems of agents that can plan, learn, and act with increasing autonomy under human oversight. These agents take ownership of entire workflows, driving outcomes with unprecedented speed and precision.

For business leaders, agentic AI moves beyond theoretical promise to deliver measurable impact. It compresses cycle times, reduces operational cost, improves service quality, and increases the speed and quality of decisions without linear headcount growth. Agentic AI will create tangible business value, directly improving organizations' profit margins. This is no longer a distant vision. According to Gartner®: "at least

15%

of day-to-day work decisions will be made autonomously through agentic AI by 2028,

up from 0% in 2024."¹ Leading organizations are already working backward from that future, redesigning workflows, decision-making processes and value-capture strategies to harness agentic AI's full potential. Yet, most enterprises are structurally unprepared to govern, secure and explain autonomous decisions at that scale.

¹ Gartner press release, Gartner predicts over 40% of agentic AI projects will be canceled by end of 2027, June 25, 2025

The untold risks of agentic AI: From hype to exposure

The countless opportunities presented by AI – it is exciting, right?

And yet, the answer is not to build more agents faster. What will set leaders apart is their ability to gain clear business value while operating agentic AI safely, securely and at scale.

As C-suite executives, you are the ones expected to deliver on someone else's narrative, and your stakeholders will hold you accountable for roadmaps, risks, and results.

Despite the hype, agentic AI is still incomplete for enterprise-grade demands. The ecosystem is racing ahead with concepts and tooling that look powerful in isolation, but brittle when exposed to real-world complexity and scale.

At a systems level, agentic AI introduces a fundamentally new class of risk.

As agents coordinate, share state and act across applications, reliability can no longer be assumed. Yet many current agentic approaches leave foundational distributed systems challenges unresolved, pushing these risks onto enterprises and integrators to absorb. In practice, this leads to familiar but dangerous failure modes: race conditions, execution deadlocks, partial completion, and non-deterministic behavior that is difficult to reproduce, explain or audit.

Without explicit orchestration, guardrails and recovery mechanisms, trust erodes, and enterprises get stuck in experimentation unable to get to production and scale.

Agentic AI fundamentally changes the threat model and expands the blast radius. Unlike traditional AI solutions, these systems have writing access, they can trigger workflows, they can and will act across systems.

A compromised agent is a digital insider. It can be manipulated through prompt injection, memory or context tampering, goal drift, tool misuse, privilege escalation, supply-chain vulnerabilities or even interactions with malicious agents. These are not theoretical edge cases, they are structural risks inherent to agentic systems, and you will be accountable for managing them. Yet despite what everyone says, organizations are applying yesterday's governance mindset to a system whose behavior changes with every context and tool call, and which requires a shift from static controls to bounded, monitored autonomy and behavioral guardrails.

A less discussed but equally structural risk is agentic sprawl. As no-code and low-code platforms make agent creation trivial, many organizations are seeing hundreds or thousands of narrowly scoped, weakly governed agents emerge through trial-and-error by end users. While democratization accelerates experimentation, uncontrolled proliferation quickly leads to duplicated logic, inconsistent behavior, unclear ownership and hidden risk. Without discipline, agentic AI risks recreating the fragmentation and shadow-application failures of past automation waves, this time at machine speed and with far greater blast radius.

The uncomfortable truth is this: the ecosystem is moving fast, but it has left you with the difficult problems pertaining to value, security, sovereignty, and scale without the industrial-grade answers you would expect.

Want to test this for yourself? Before signing off on another agentic AI initiative, try asking your teams these six simple questions and pay close attention to how confidently, or vaguely, they answer. The questions expose whether agentic AI is being treated as an experiment or ready for production-grade.

1.

Identity, access and control: Do we know under which identities our agents run, what permissions they inherit, how privilege escalation is prevented, and how access can be revoked in real time when agents call tools and APIs autonomously?

2.

Failure propagation: If a prompt, memory, or input is manipulated, do we understand how that vulnerability propagates across an end-to-end workflow and across multiple agents, rather than assuming it is a contained model issue?

3.

Traceability and auditability: Can we trace, audit, and explain every action an agent takes in a way that stands up to regulatory, security, and operational scrutiny?

4.

Data governance and trust: Do agents operate on curated, reliable, and governed data sets with clear ownership and quality controls, or are we exposing autonomous systems to inconsistent, stale, or untrusted data?

5.

Orchestration and discipline: Do we have a clear orchestration layer that manages agent coordination, state, handoffs, failure modes, and escalation, or are we relying on brittle point-to-point integrations?

6.

Cost control and lock-in risk: Do we have mechanisms to monitor, control, and optimize probabilistic and highly variable agentic AI costs (tokens, tools, retries, loops) and avoid future lock-in, or are we assuming spend and dependencies will remain manageable as autonomy scales?

These risks are real, but they are solvable.

With the right foundations, enterprises can scale autonomy with confidence

Atos's expertise lies in solving complex, high-stake problems in the toughest environments enterprises operate in, across public and private sectors alike. This includes brownfield realities, heavily regulated sectors, mission-critical operations, sovereign infrastructures, and ecosystems under constant cyber scrutiny.

This includes environments held to military-grade standards, where reliability, high security, resilience, and traceability are non-negotiable, and where systems are designed to operate under attack, failure, and extreme operational constraints. These are contexts where software must work, fail safely, preserve command and control, and stand up to continuous audit and regulatory oversight.

Why sovereignty matters in agentic AI

Sovereign AI means deliberate control over critical data, decisions, and AI behavior. Atos enables enterprises to choose where sovereignty truly matters and where commoditized ecosystems suffice. We help to ensure you never have to choose between frontier-level capability and independence from another provider's roadmap or jurisdiction. Through sovereign-ready infrastructure, flexible deployment options (air-gapped, private cloud, hybrid) and partnerships, Atos delivers autonomy without compromise.

Atos brings a rare combination of governance discipline, orchestration expertise and experience operating mission-critical and sovereign systems. We treat agentic AI as production software that is governed, secured and operated to help organizations scale autonomy deliberately while preserving control, auditability, and trust. Critically, Atos leads with autonomy only when governance, security and measurable value are proven.

- ✓ We operate agentic AI with sovereign governance including fine-grained control over data quality and residency, model deployment, access boundaries and operational oversight where required.
- ✓ Our delivery and commercialization model is focused on outcomes driven by value tracking, KPI alignment and cost governance tied to agent performances.
- ✓ We operate a production-grade agentic management platform, underpinned by deep engineering expertise, that ensures predictable agent behavior, enforceable cost guardrails, policy-controlled tool usage and resilient execution at enterprise scale,
- ✓ We mobilize cross-IT capabilities, assets, and tools to address the full stack, from infrastructure, platforms to applications, operations and cyber.
- ✓ We make security a competitive advantage with runtime behavioral security, privilege boundaries, real-time revocation, context-poisoning defenses and continuous auditability.
- ✓ We leverage a deep partner ecosystem, where we act as the orchestrator, combining hyperscalers, enterprise platforms, and specialized players while preserving architectural independence, transparency and freedom of choice for you.
- ✓ We anchor agentic orchestration in a human-centric operating model, embedding change, accountability and ethical guardrails into how work is actually done.

Today, you may have probably tested a few agents or may already have pilots running somewhere in your organization. The real question is whether you trust them to scale. **What differentiates outcomes is not ambition, but execution discipline.**

This is where Atos comes in. We have spent decades operating complex, regulated, and mission-critical systems, and agentic AI is no exception. Collaborate with us to build the foundations on which you can scale autonomy with control, trust, and measurable results.

Read on to understand how this works in practice and how agentic AI can move from promise to production.

Our vision:

Redefining services for an agentic world

Atos has a clear and unapologetic vision of the future: Services as Software (SaS)

We believe enterprise services are converging with enterprise technology into a new operating layer where routine, repeatable and data-intensive work is delivered as software, and less as manual people-based services. In this model, value moves from selling efforts to shipping capabilities; from seat-based (time and material) deliveries to outcome-based products. This is a structural reset of how services are designed, delivered and priced.

Atos is already moving decisively in this direction. We are investing in agentic driven delivery for autonomous operations, accelerated modernization, software engineering and cyber-defense through an end-to-end agentic platform strategy to productize what were historically people-led services. At the same time, we are redesigning our own delivery and G&A processes as case studies for service-as-software models. This includes rethinking how work is distributed into workflows, how agents and humans collaborate, how outcomes are measured and how continuous improvement is shipped through software releases.

Atos is not asking you to go somewhere we have not gone ourselves. Atos is already applying this model internally, by acting as our own Client Zero. We are building practical experience in agent orchestration, cost control, security and governance, and using years of experience in complex environments through anonymized data. These lessons are crucial inputs that enable us to decide how we can help clients industrialize agentic AI in real environments.

Atos is redefining how agentic AI is operated in enterprises and public institutions.

Agentic systems leverage insights to drive actions and a growing share of operational and strategic decisions across your organization. Atos provides the foundations to keep those decisions controllable, auditable, secure and sovereign, where required.

Our ambition is backed by a trusted control plane for agentic AI across the most demanding environments, ensuring decision-making at scale remains transparent, accountable, and compliant with regulatory and institutional expectations. As you increase your autonomy, we will ensure you retain control, traceability and trust in how decisions are made and executed.



Responsible AI: A condition for scalable autonomy

As agentic systems take on a growing share of operational and decision-making responsibilities, Responsible AI (RAI) becomes a core operating principle of how services are designed and run. Agentic AI will be governed far less by AI-specific regulations than by industry, process, and functional rules that already apply to enterprise operations (i.e., privacy, security, financial controls, HR, safety and trade). Responsible AI therefore starts with the fundamentals like security, privacy, fairness, accuracy, explainability, but must extend to how autonomy is explicitly bounded, supervised and updated as business context, regulations and expectations evolve.

In an agentic world, Responsible AI and governance converge. Responsibility means agentic systems are aligned with business strategy and corporate values, monitored at process level, and designed with deliberate human-in-the-loop and system-level controls. Humans must be able to approve, override, interrupt, or suspend agents and entire workflows when risk or impact requires it. This is how autonomy remains compatible with trust at scale. Agentic AI only becomes a durable advantage when it can be observed, controlled, and held accountable in day-to-day operations.

Trust at scale with secure agentic AI

Security and AI governance sit at the heart of our agentic AI strategy, because we believe autonomy without control is an unmanaged risk. As agentic systems begin to execute work across your applications and processes, trust must be designed into how they operate: who is accountable, what is permitted, how behavior is supervised and how actions are verified. This is why Atos treats security and governance as part of the operating layer we deliver instead of external checks applied after the fact.

Our approach is to make trust scalable. Governance provides the lifecycle discipline of ownership, risk-based guardrails, human oversight, change control, trustworthy assurance so that agents remain aligned as workflows evolve and capabilities expand. Security provides the enforcement discipline to make sure access, actions, and data usage remain controlled in live environments. Together governance and security allow to increase autonomy progressively while keeping execution transparent, auditable, and defensible under regulatory, cyber and sovereign constraints.

Sovereign Agentic AI

Control where autonomy matters



Sovereign AI does not mean technological isolation or rebuilding the AI stack from scratch. It is about deliberate control over what matters most, and explicit choices about where sovereignty is required versus where commoditization is acceptable. Most organizations will continue to rely on external models, platforms, and ecosystems for non-differentiating use cases. Sovereignty becomes critical when AI systems drive core business decisions, act autonomously across workflows, handle sensitive or regulated data, or encode proprietary knowledge and processes. The question is therefore not “Are we sovereign?”, but “Where do we need sovereignty, and to what degree?”

In an agentic context, sovereignty cannot be addressed at a single layer. If agents can take actions, call tools, interact with other agents, and modify systems, any weak link becomes a systemic vulnerability. Sovereign Agentic AI must therefore be enforced across the full stack: data, models, orchestration layers, applications, and runtime operations. This includes governance over how agents behave and evolve, transparency and traceability of decisions and data provenance, predictable and robust model behavior, and end-to-end control and security of data in execution.

Sovereignty is not a static architecture choice but an operational discipline, embedded into how agentic systems are monitored, constrained and secured in real time. This is the only way autonomy remains resilient, auditable and controllable at scale.



Vision to reality

A disciplined path to autonomy

You may operate mission-critical environments under regulatory, industry, sovereign and cyber constraints where failure is not an option, but a transition towards an agentic-first model cannot be rushed or imposed. Scaling our vision will therefore be progressive, disciplined, and responsible.

It is our mission to help you move deliberately from today's human-centric services to hybrid models, and eventually to services as software-delivered operations without breaking what already works. We will scale autonomy where it creates value, keep humans in the loop where needed, and industrialize only what can be secured, governed, and measured.

Just as importantly, we are pragmatic about where agentic AI truly adds value. Not every workflow requires agentic orchestration. For highly deterministic, standardized processes, simpler automation approaches can deliver faster ROI with lower risk and overhead. Our role is to help you make the right trade-offs, deploying agentic AI where adaptability and judgment matter, and proven automation where predictability and efficiency win.

The implication is that not all services should evolve in the same way, or at the same pace. The right balance between human judgment and software-driven execution depends on two factors: the complexity of the work and the level of human empathy and trust it requires. **The Service Value Matrix below illustrates how different types of services naturally lend themselves to different delivery models, from fully automated efficiency machines to human-led, AI-assisted advisory roles.**

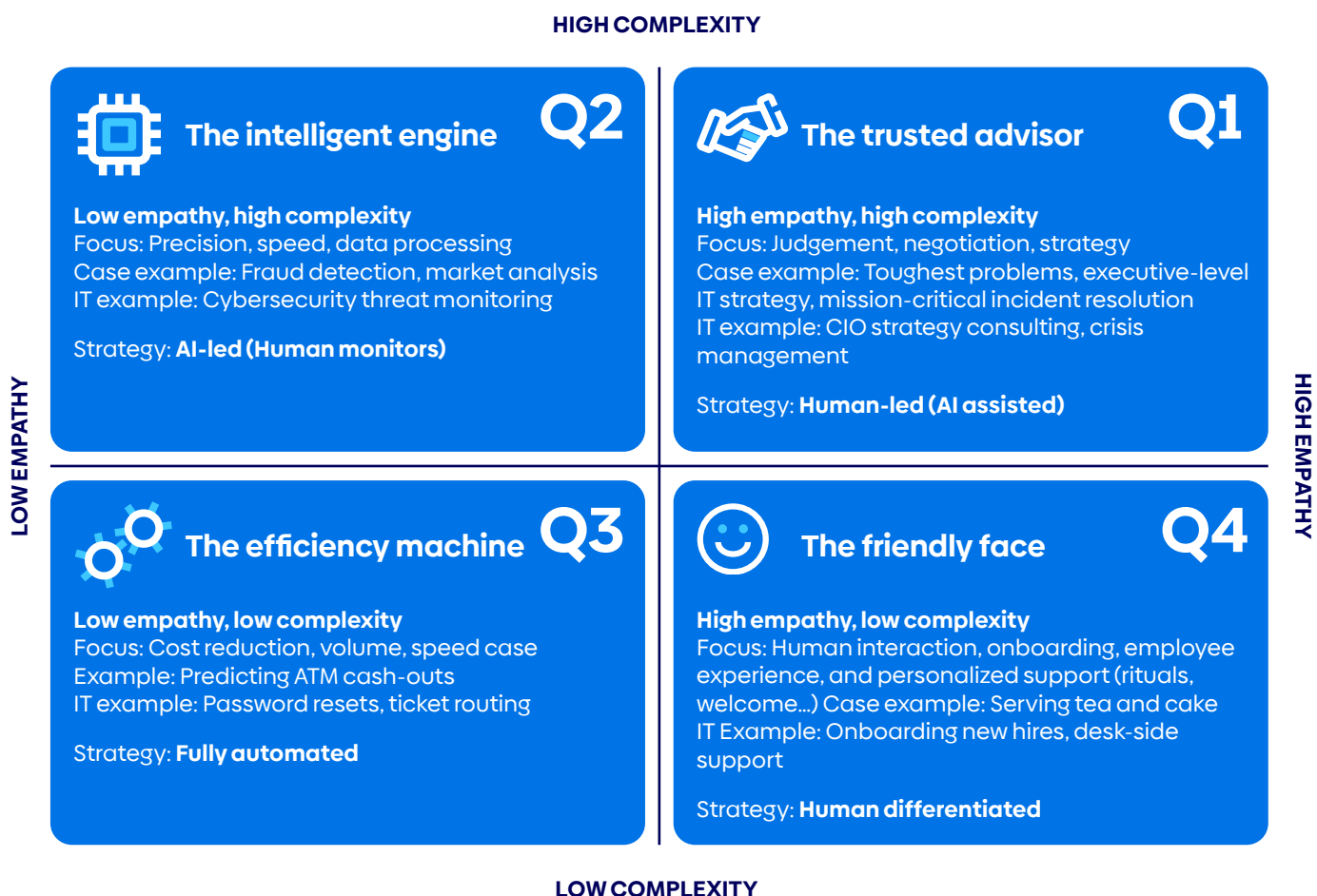


Figure 1. Where agentic AI actually makes sense, the service value matrix

In practice, becoming AI-native follows a progressive maturity path. Your organization will move from human-led execution, where AI assists analysis and decision-making, to agent-assisted models where agents execute defined tasks under human supervision. From there, agent-orchestrated workflows emerge, with humans managing exceptions and oversight. Only then does bounded autonomy become viable, where agents own end-to-end workflows within clearly defined controls.

Moving along this maturity path is more about choices than it is about technology. Each step requires deliberate decisions on ownership, governance, architecture, and operating model. In practice, this is where your organization might be at crossroads. Industrializing agentic AI requires both discipline and ambition. Based on what we see in client environments and early large-scale deployments, a small set of choices separates organizations that move into production from those that remain stuck in experimentation.



The following dos and don'ts set clear stakes in the ground for scaling agentic AI responsibly and delivering real outcomes.

Do's

- ✓ Do make this a CxO/ Board-level outcome program with named business owners and clearly defined ROI objectives.
- ✓ Do choose a small number of end-to-end workflows with clear value, repeatability and high ROI potential.
- ✓ Do baseline KPIs upfront and commit to measurable targets in production environments.
- ✓ Do put AgentFinOps in place early, with unit economics, spend visibility, anomaly detection, and throttling as autonomy scales.
- ✓ Do invest in modern data platforms and strong governance to give agents reliable, secure, high-quality data foundations.
- ✓ Do redesign AI-first workflows, including decisions, exceptions, handoffs, and agent-human interactions.
- ✓ Do assign agents explicit scopes, decision rights, escalation thresholds and human supervision models.
- ✓ Do establish robust communication and orchestration mechanisms (APIs, messaging, and shared state) between agents and systems.
- ✓ Do implement AgentOps from day one, with observability, audit trails, rollback, kill switches, and safety controls.
- ✓ Do scale progressively and deliberately: pilot fast, iterate faster, and expand autonomy based on evidence and learning.
- ✓ Do secure agentic-specific risks through runtime, behavioral security backed by access control, encryption, privacy safeguards and egress controls, with defenses against context poisoning, goal drift, tool misuse and privilege escalations.

Don'ts

- ✗ Don't run it as IT pilot initiatives with business teams watching and agents operating without ownership.
- ✗ Don't start with scattered use cases or low-impact experiments disconnected from P&L outcomes.
- ✗ Don't rely on demo metrics, anecdotal productivity claims, or pilot-only success.
- ✗ Don't let token, tool, or loop-driven costs grow unchecked.
- ✗ Don't deploy agentic AI on fragmented, poorly governed data ecosystems and expect consistent, trustworthy outcomes.
- ✗ Don't automate broken processes or layer agents onto existing inefficiencies.
- ✗ Don't allow uncontrolled autonomy or ambiguous accountability for outcomes and failures.
- ✗ Don't rely on ad-hoc interactions that create coordination failures, duplication, or silent error.
- ✗ Don't treat governance, safety, and compliance as phase-two topics.
- ✗ Don't jump to hands-off execution or ignore the change management required to build a learning organization.
- ✗ Don't assume generic GenAI guardrails are sufficient for autonomous, write-enabled agents.

Enterprise-grade governance, security, and sovereignty

Answering difficult questions about agentic AI: A disciplined path to autonomy

Atos articulates its vision for agentic as follows. Agentic AI will not succeed as thin wrappers around models. It must be engineered and operated as production software running inside production systems. The questions raised earlier in this document are actual, concrete failure modes that organizations encounter when moving from experimentation to live operations. Atos's Agentic AI offerings are designed to address these challenges directly, through architecture, operating discipline, and enforceable controls.



Identity, access, control, and human oversight

Agents operate under explicit non-shared identities with tightly defined permissions. Access to tools and systems is governed through zero trust principles, enforced centrally, with just in time access, where appropriate, and the ability to revoke access in real time. Privilege escalation paths are controlled and auditable, ensuring agents cannot silently expand their authority as autonomy increases and that high impact actions can be constrained to the right identities and approval paths. Agent goals, intent, and boundaries are explicitly defined and continuously monitored, with mandatory human oversight for goal changes and confidence-based escalation to prevent a silent drift away from the original business intent.



Data quality and sovereignty by design

Agents operate on curated, governed, and trusted data sets with clear ownership and lifecycle management. Data readiness, provenance, sovereignty constraints, and privacy requirements are assessed upfront and enforced continuously, ensuring autonomy does not amplify inconsistent, stale, or untrusted information.



Orchestration and architecture at scale

Atos provides a dedicated orchestration layer and reference architectures including secure zero trust multi context protocol (MCP) servers to manage coordination between agents, systems, and humans. State awareness, handoffs, escalation logic, and recovery mechanisms are engineered explicitly, including deterministic rollback paths, kill-switch authority, and controlled restart of agents and workflows, avoiding brittle point-to-point integrations and enabling predictable behavior at scale.



Security propagation across workflows with secure agentic supply chain

Atos designs agentic workflows as distributed systems, with explicit orchestration, state management, and failure handling. Security controls such as prompt-injection defenses, input validation, memory isolation, and tool-use constraints, are applied at workflow level and not at model level. This is across the full chain of inputs, retrieval, memory, coordination, and tool execution. Architectural mechanisms are in place to isolate failures and prevent vulnerabilities from propagating across agents and systems. This includes securing the agentic supply chain itself: third-party agents, partner components, model dependencies, and external tools are vetted, versioned, and governed as production dependencies, with continuous integrity checks to prevent compromised components from entering live workflows.



Traceability, auditability, accountability, and response readiness

Every agent action is observable and traceable. Decisions, tool invocations, data access, and outcomes are logged in a way that supports operational analysis, regulatory reporting, and security investigation. This traceability underpins response readiness, enabling rapid investigation, controlled rollback, and decisive intervention when anomalies, policy breaches, or failures are detected.



Cost discipline and value-linked scaling

AgentOps and FinOps are implemented from day one, tying agent run-cost to business KPIs like cycle time, cost-to-serve, throughput, quality, and tracking correlation as autonomy expands. Real-time monitoring and controls are in place before scale, enabling teams to allocate budgets and throttle or stop spend at the right level (feature, workflow, project, program, business unit). Reusable, governed agents are curated and shared through a controlled library, promoting reuse over duplication and preventing agent sprawl while accelerating time to value. This allows teams to manage token usage, tool calls, retries, and execution loops while preserving freedom of choice and limiting lock-in.

Atos's Agentic AI approach

From blueprint to industrial scale with Atos sovereign agentic studios

Atos's approach to agentic AI is grounded in a simple conviction: scaling agentic AI is more than a technology challenge. The business, operating model and human factors are equally important. Successful deployments follow a clear blueprint that we have observed in leading organizations and operationalized through our own offering and internal transformation.

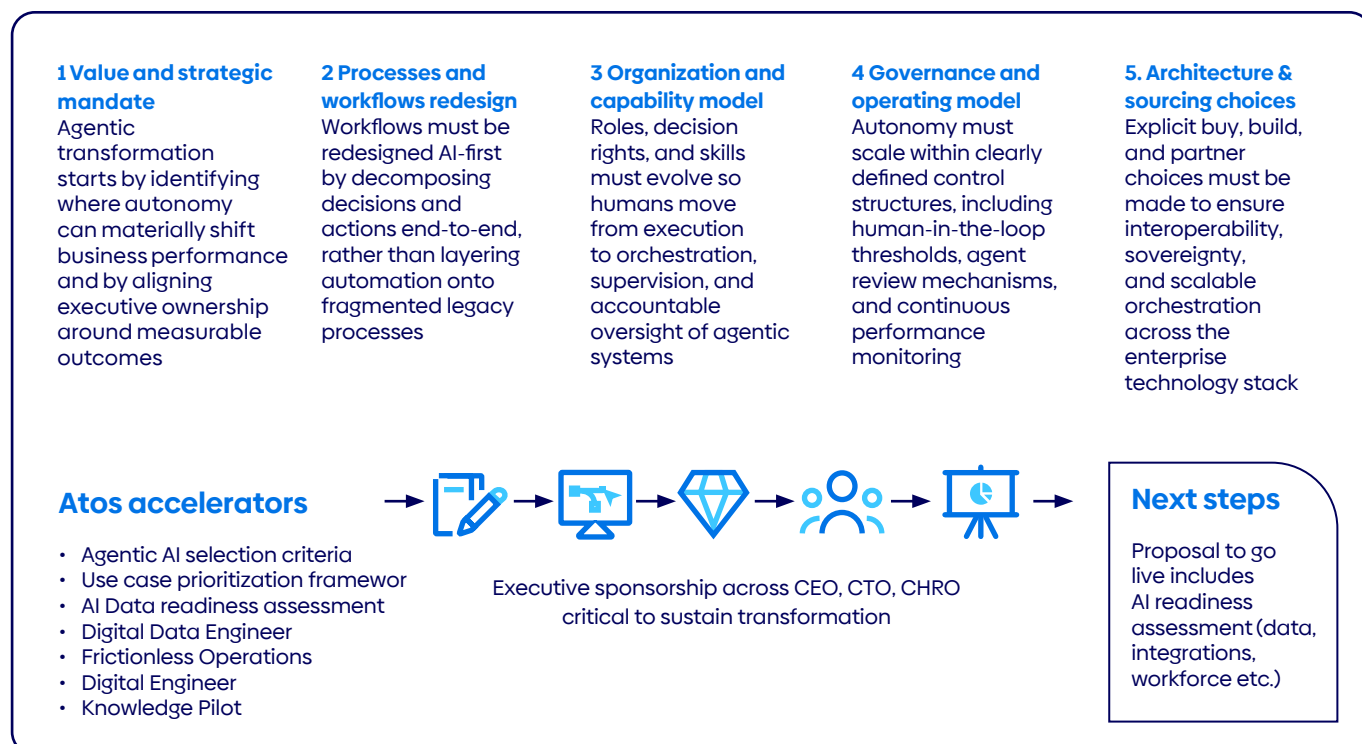
In practice, this blueprint is implemented through Atos Sovereign Agentic Studio. It is a production-grade capability designed to engineer, run, and govern agentic systems in live enterprise environments and accelerate time-to-value from agentic AI through a layered architecture. Take a look:

Business interaction and value delivery at the top, orchestration and governance at the core, agent cognition and data context underneath, and sovereign-ready infrastructure as the foundation. **This is how Atos moves organizations from human-centric operations to AI-native delivery.**



Additionally, we require a disciplined, end-to-end approach to transforming how work gets done. We apply a clear agentic blueprint that starts from concrete business value and priority workflows, redesigns workflows AI-first, establishes trusted and secure foundations, and then scales agentic execution responsibly.

The figure below summarizes how Atos helps organizations move from human-centric operations to AI-native delivery.



How Atos delivers this in practice

Atos's Agentic AI offering aims to be a production capability to design, deploy, and operate autonomous agents that executes end-to-end business and IT processes. It is built on existing Atos delivery assets, including our digital engineer powered managed services, digital transformation engineer powered modernization capabilities, industry solutions, and the Atos Polaris acceleration platform. It is designed to run agents in live environments, including regulated, sovereign, and mission-critical contexts. It runs on the following main nine pillars:

1. Agentic launchpad

Atos begins with a comprehensive business and technology audit to assess process maturity, data and application landscape, integration readiness, sovereignty constraints, and cyber exposure. This includes an Agentic Readiness Assessment, combining Atos's own readiness framework with selected partner tools to assess preparedness across people, processes, data, architecture, security, and operating model. We then run an AI-first workflow redesign to identify and prepare the highest-value agentic opportunities, define measurable success criteria, and build a prioritized roadmap from controlled pilots to scaled production. This front-end diagnostic ensures that when we deploy the Agentic AI Studio, it is anchored in value creation from Day 1 and lands in an environment that is governable, secure, and set up to deliver real outcomes.

2. Sovereign-ready build-and-run environment

This is a sovereign-ready environment to design, test, and operate agentic workflows across public, hybrid, private, and sovereign cloud/AI infrastructure setups. It combines low-code tooling for process owners, curated connectors into existing IT (ERP, CRM, OT, ticketing), and secure access to multiple model and tooling options to avoid single-stack dependency. When sensitivity demands it, we can active / activate? privacy-preserving encryption tools so sensitive data and workloads remain protected throughout processing. And where required, we can deploy in sovereign/on-prem configurations aligned with data residency and compliance constraints.

3. Agent orchestration and reference architecture

Atos designs and implements the orchestration layer via its Digital Engineer autonomous agentic operations framework that allows multiple agents, systems, and humans to work together, reliably and at scale. This includes defining reference architectures for agentic systems, coordination patterns between agents, state and context management, error handling, escalation logic, and human-in-the-loop handoffs.

These architectures are grounded in a Reference Agent Framework that standardizes agent roles, interaction patterns, state handling, and termination conditions, and leverages MCP-compliant communication models to ensure interoperability and predictable behavior across agents and tools. The orchestration layer ensures agents do not operate as isolated components, but as part of a controlled end-to-end system with predictable behavior, clear ownership, and deterministic failure modes. The Atos Polaris AI platform delivers the essential foundational components required to implement the reference architecture for secure, scalable and reliable agentic orchestration, including shared context services, execution state tracking, confidence-based escalation mechanisms, and controlled retry and rollback patterns. It offers a pre-integrated suite of services that capitalize on the native capabilities of leading hyperscalers and incorporate best-in-class open-source technologies.

4. Security and governance control plane

Atos provides a unified security and governance control plane for agentic systems in production. This control plane enforces security and policy constraints at runtime, defining what agents are allowed to do, under which conditions, and with which escalation and approval paths. It includes continuous behavioral monitoring, auditable end-to-end traceability of agent actions and decisions, and lifecycle governance covering onboarding, updates, versioning, and retirement of agents, models, tools, and connectors. This governance is embedded into the delivery model through concrete operating frameworks, including defined agentic roles (e.g. AgentOps Lead, AI Workflow Architect), a formal Agent Review Board governing changes in agent behavior and autonomy, and standardized release, validation, and rollback procedures aligned with the Agentic AI lifecycle. Reference agent and orchestration frameworks, including MCP-compliant interaction patterns, are used to enforce consistent security boundaries, state management, and escalation logic across workflows. By combining runtime controls, immutable audit trails, and disciplined release and validation gates, Atos ensures agentic autonomy remains aligned with business intent, regulatory requirements, and risk posture as systems evolve over time.

5. MCP integration services

Atos strategically delivers design, implementation, integration, and operate at scale Model Context Protocol (MCP) services. MCP Integration Services is an end-to-end implementation capability that bridges AI, cloud platforms, data, and enterprise systems, ensuring MCP services are production-ready, governed, and resilient from day one to provide secure, controlled and governed end-to-end agent communications.

6. AI-backed application modernization

Atos delivers Application Modernization and engineering at scale, from partially automated to seamless, autonomous and intelligent application transformation service delivery. The Digital Transformation Engineer is an Agentic AI-powered platform that deploys autonomous, goal-oriented AI agents to simplify and accelerate how we assess and deliver application transformation services at scale. It combines the precision of Agentic AI with the assurance of human oversight supporting the evolution to AI-SDLC.

7. AI Agents marketplace

Clients draw from a curated library of outcome-backed agents, pre-built by Atos, sourced from trusted partners, or clients themselves, spanning horizontal and industry-specific domains. Each agent ships with a defined scope, guardrails, integration patterns, and measurable KPIs to accelerate time-to-value without sacrificing control. This composable portfolio balances bespoke builds for differentiating workflows with reusable agents at scale, enabling multi-vendor orchestration across hyperscalers and specialist players, while preserving sovereignty, accountability and freedom of choice.

8. Process redesign, AgentOps, and change management

Atos co-designs AI-supervised workflows with business, IT, risk, HR, and security teams, then industrializes them through AgentOps. Governance, observability, cost control, and policy management are embedded from day one, alongside workforce enablement and AI-first skill development. This is what allows clients to move progressively from human-led execution to agent-orchestrated execution, with humans supervising where required, and expanding autonomy only when it is secure, auditable and measurable.

9. Deep partner ecosystem

Our delivery model is underpinned by a broad and curated partner ecosystem. Atos provides platform-neutral orchestration across hyperscalers, model providers, and enterprise systems, ensuring true freedom of choice and avoiding lock-in at both the orchestration and infrastructure layers. Operating at the intersection of hyperscalers, core enterprise platforms, and emerging agentic technologies, Atos combines access to large-scale compute and AI infrastructure with best-of-breed agentic components. We support public, hybrid, private, and sovereign environments by leveraging hyperscalers for scale and performance, enterprise platforms to embed agents directly into business workflows, and specialized partners for agent frameworks, orchestration accelerators, and operations at scale. This ecosystem-driven approach allows Atos to avoid single-vendor dependency, assemble the right stack for each context, and maintain a unified control plane for governance, security, cost management and sovereignty.

Beyond technology, the sovereign agentic studio operates as a co-innovation and orchestration engine. It brings together cutting-edge applied research, selected start-ups and emerging agentic frameworks, and enterprise-grade platforms, allowing Atos to rapidly test, assemble, and industrialize agentic solutions that are ready to scale. Clients benefit from innovation without bearing the cost, risk, and fragmentation of navigating the ecosystem alone.

4. Agentic AI already at work

An Atos case study

Across operations, engineering, and transformation programs, Atos has moved beyond experimentation to deploy Agentic systems that take ownership of real workflows under security, governance, and human oversight constraints. The following examples illustrate how we are reshaping execution today, and how organizations can progress from assist-ed automation to orchestrated autonomy in a controlled way.

CASE STUDY

Scaling incident resolution under SLA pressure with Agentic AI Digital Engineers (DE/ DTE)

Business challenges:

In large, business-critical IT environments, incident management has become a structural performance bottleneck. As release frequency increases and system landscapes grow more complex, engineering teams are overwhelmed by rising alert volumes, manual triage, and fragmented diagnostic processes spread across multiple tools. Root cause analysis depends heavily on individual expertise and institutional knowledge, making outcomes inconsistent and difficult to scale. In this context, the Mean Time To Repair typically ranges from three to four hours, SLA breaches become more frequent, and operational costs continue to rise as organizations attempt to compensate with additional headcount.

Atos's solution:

Atos worked with the client to fundamentally change how incident handling was executed. Rather than using AI to assist engineers at the margins, the client introduced agentic execution to take ownership of the incident resolution flow itself, under strict security, governance, and confidence-based controls.

Agentic systems were deployed to continuously monitor observability data, correlate alerts across applications and infrastructure layers, analyze recent deployments and service dependencies, and diagnose likely root causes using historical incident patterns and operational runbooks. When predefined confidence thresholds were met, agents automatically executed approved remediation actions such as restarts, scaling operations, rollbacks, or feature toggles. When uncertainty remained, incidents were escalated to human engineers with full diagnostic context and clear recommendations, allowing humans to focus on judgment and exceptions rather than repetitive investigation. This shift from human-led firefighting to agent-owned execution delivered measurable results in production:

20-40%

reduction in Mean Time To Repair (MTTR), driven by faster diagnosis and automated remediation

55-75%

of incidents resolved end-to-end without human intervention, under confidence-based controls

40-60%

reduction in L1/ L2 ticket volumes, significantly lowering operational noise and alert fatigue

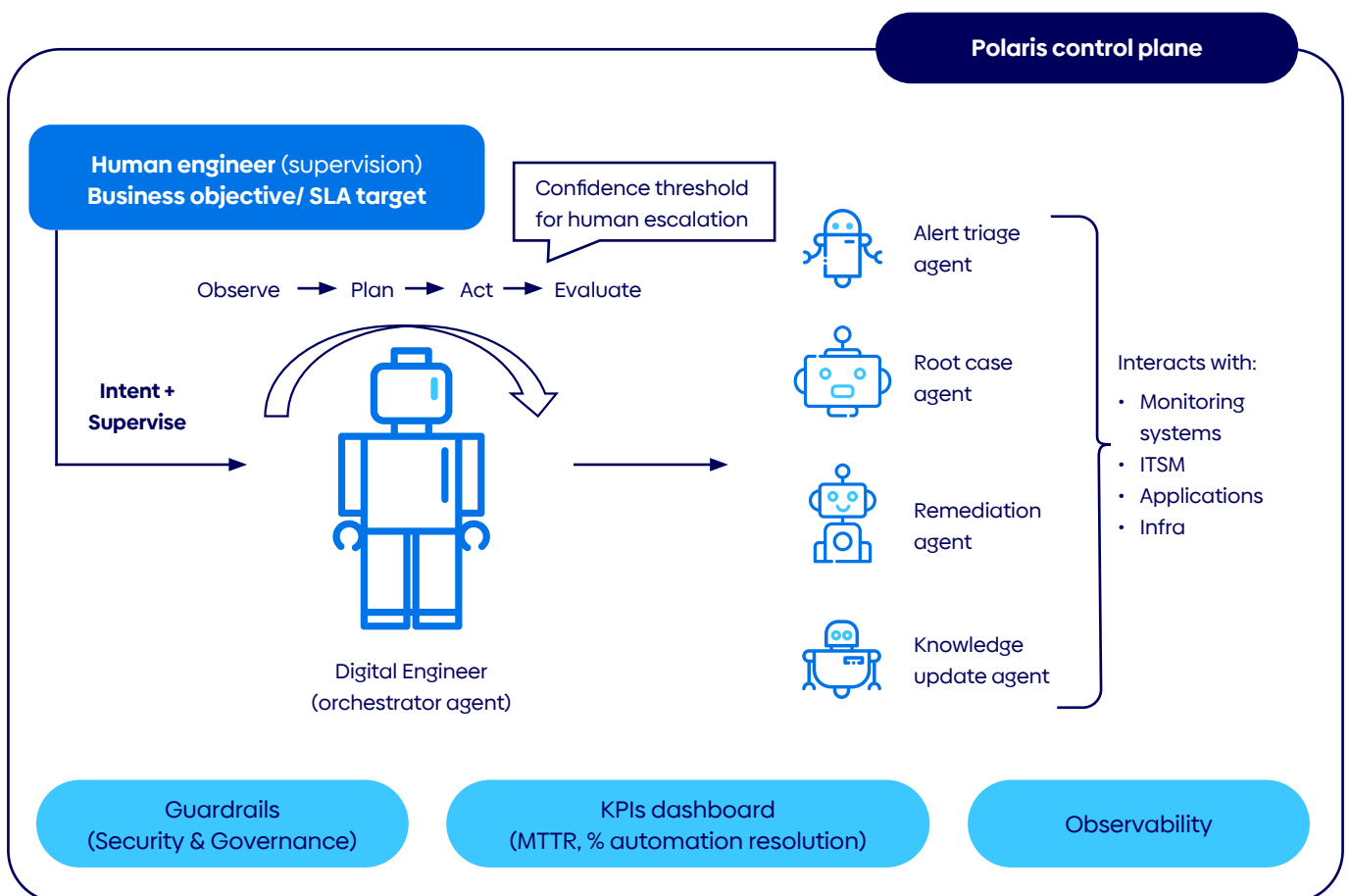
20-35%

reduction in operational costs, while improving service reliability and SLA adherence

Beyond the numbers, the operating model itself changed. Engineers moved away from constant reactive work toward supervision, optimization, and resilience engineering. Incident handling became more predictable, auditable, and easier to operate at scale, even under sustained SLA pressure.

Crucially, autonomy was introduced progressively. All agent actions were logged, auditable, and reversible. Service-level objectives were continuously monitored, with automatic rollback or escalation when recovery was unsuccessful. ITSM tickets, post-incident reviews, and knowledge base updates were generated automatically, ensuring full traceability and compliance.

What made this transformation possible was not a single agent, but an industrial execution backbone. Behind the scenes, Atos leveraged its Digital Engineer and **Digital Transformation Engineer capabilities**, powered by the Atos Polaris platform, to orchestrate and operate agentic workflows reliably in production. Polaris provided the control plane that enforced security, observability, cost discipline, and confidence thresholds, allowing agentic execution to scale safely. The client did not adopt a new toolset; they benefited from a new way of executing work, where autonomy increased without sacrificing control.



Glossary

APIs: Interfaces that let software systems exchange data and trigger actions

Agile: An iterative delivery approach based on short cycles, frequent feedback, and continuous improvement

FinOps: Practices to forecast, track, and optimize the cost of running AI agents (tokens, tools, infra)

Agent orchestration: Coordinating multiple agents, tools, and systems to execute a workflow reliably end to end

Agentic AI: AI that uses agents to plan, decide, and act via tools to complete tasks with varying levels of autonomy

AgentOps: Operating model to deploy, monitor, govern, and improve agents safely at scale

Autonomous agents: Software agents that can execute actions with minimal human intervention within set guardrails

Brownfield (software development): Building or modernizing on top of an existing, complex legacy environment

CRM: Customer Relationship Management system used to manage sales, accounts, and customer interactions

DevOps: Practices and tooling that combine software delivery and operations to release reliably and continuously

ERP: Enterprise Resource Planning system that runs core processes like finance, procurement, and supply chain

G&A: General and Administrative overhead costs such as finance, HR, legal, and corporate functions

GenAI: Generative AI that produces content (text, code, images) from prompts and context

Hyperscalers: Large cloud providers that offer elastic compute/storage at global scale - like AWS, Microsoft Azure, and Google Cloud Platform

Infrastructure stack: The layers that run applications and compute, storage, network, platforms, and runtime services

Kill switches: Controls to immediately stop an agent or workflow when risk, drift, or failure is detected

LLM: Large Language Model is a type of AI model trained on massive amounts of text (and sometimes code, images, audio) to predict and generate language. In practice, it can understand prompts, answer questions, summarize, write, translate, and reason through tasks by generating the next most likely tokens in a sequence

MCP: Model Context Protocol is an open protocol that standardizes how an AI app (an “LLM client/host”) connects to external tools and data sources through MCP servers, so the model can fetch live context and do things like run a tool, query a system, retrieve a file/resource in a consistent way

Mission-critical contexts: Systems and operations that must run reliably with strict resilience and compliance requirements

MTTR: Mean Time To Restore/Repair, average time to recover service after an incident

Orchestration layer: The runtime that routes tasks, manages state, enforces policies, and integrates tools/systems

OT: Operational Technology, industrial/control systems that monitor or control physical equipment and processes

Reference agents: Standardized, reusable agent patterns/templates with defined scope, controls, and KPIs

Service as a software: Delivering an operational service primarily through software automation rather than manual efforts

Sovereignty: Requirements that data, models, and operations stay under specific jurisdictional or national control

Token: A unit of text processed by an LLM, used to measure usage and cost



Let's talk!

As Agentic AI accelerates, the organizations that win will pair autonomy with control, treating agentic systems as production software with governance-first architecture, zero-trust access, behavioral security, auditability, and cost discipline from day one.

Atos Sovereign Agentic Studio provides the industrial backbone to do this in regulated and mission-critical environments, preserving sovereignty, transparency, and freedom of choice while delivering measurable outcomes.

Engage with us! We will help you develop one high-value, end-to-end workflow - confirm data readiness, define decision rights and human-in-the-loop paths, instrument AgentOps/FinOps - and run a production-grade pilot with observability, rollback, and kill switches—then scale what works.

Next Step: Connect with Atos to select a value stream and stand up the control plane to scale, on your terms.

Connect with us at agenticai@atos.ai

Atos Group's tech ambition

Atos Group's strategy is built around three closely interconnected priorities: agentic AI for mission-critical environments, digital sovereignty, and cybersecurity. Together, they enable organizations to innovate with confidence, retain control of critical digital assets, and operate securely in an increasingly complex world. This ambition is brought to life through the complementary strengths of Atos, Atos Amplify and Eviden.

Atos

Atos is the services brand of Atos Group, delivering AI-powered, end-to-end digital services and solutions to public and private organizations worldwide.

Atos designs, builds and operates digital environments that are critical to performance, resilience and sovereignty, helping clients maintain control over their data, infrastructure and operations while accelerating transformation. From cloud, data and AI to cybersecurity, applications, digital workplace and managed services, Atos supports organizations throughout their entire digital journey.

Trusted to run complex, business-critical environments, Atos helps clients modernize core IT, strengthen security, improve operational efficiency and create sustainable business value in an increasingly digital world.

atos.ai

Atos Amplify

Atos Amplify is Atos Group's AI-ready consulting and advisory business, supporting organizations make critical strategic decisions and turn them into measurable business outcomes.

Combining deep industry expertise, business insight and AI-driven innovation, Atos Amplify works with clients to define transformation strategies, accelerate adoption of emerging technologies and unlock value from data, AI and digital platforms. Its consult-to-build approach connects strategy with execution, helping organizations navigate complexity, strengthen sovereignty, improve resilience and drive sustainable growth.

Grounded in strong human values and responsible AI principles, Atos Amplify enables organizations to transform with confidence, transparency and accountability.

atos.ai/atos-amplify

EVIDEN

Eviden is the Atos Group brand focused on cybersecurity products, mission-critical systems and vision AI. A trusted partner to public and private sector organizations, Eviden deploys advanced analytics capabilities and helps ensure the robustness, continuity and resilience of its customers' operations.

Its portfolio combines advanced cybersecurity technologies, digital identity solutions, data protection capabilities, and mission-critical systems designed for the most demanding environments along with real-time AI-powered monitoring capabilities with AI.

Leveraging more than 2,200 experts and 720 patents, Eviden helps protect people, data and critical infrastructures worldwide, at the convergence of intelligence, sovereignty and trust.

eviden.com

About Atos Group

[Atos Group](#) is a global leader in digital transformation with c. 56,000 employees and annual revenue of c. €7.2 billion (at the go-forward perimeter), operating in 54 countries under two brands - Atos for services and Eviden for products and systems. European number one in cybersecurity and a leader in cloud, Atos Group is committed to a secure and decarbonized future and provides tailored AI-powered, end-to-end solutions for all industries. Atos Group is listed on Euronext Paris.

Find out more about us
atos.ai

Let's start a discussion together

