



From Pledges to Practice:

Atlassian's Commitment to the EU AI Pact



Atlassian is an enterprise software company guided by a singular mission: to unleash the potential of every team. Our software helps teams do their best work across software development, service management, and knowledge and work management. If you've ever filed a Jira ticket, collaborated on a Confluence page, or made a Trello board, then you've used an Atlassian product.

We are proud to support European economic growth and enterprise productivity. Teams from across 85% of the Fortune 500 and over 300,000 companies worldwide – including Air France-KLM, Riverty, Amadeus, Software AG, and Mercedes-Benz – rely on our solutions. Roughly one third of our customers are based in Europe, who range from global enterprises to growing startups. For example, [Glovo](#), an on-demand delivery platform based in Spain partnered with Atlassian to optimize their business processes and enhance internal user experience.

Our commitment to the [EU AI Pact](#) is rooted in our Responsible Technology Principles. In 2023, we published our [Responsible Technology Principles](#), which reflect our company values and steer our approach to responsible AI. To operationalize our principles, we developed our own [Responsible Technology Review Template](#) and published our [No BS Guide to Responsible AI Governance](#) to describe our learnings from conducting dozens of reviews across Atlassian. Our commitment to the Pact is another step in our journey of translating our Principles – and our Pact pledges – into practices.

Atlassian adopted five pledges under the Pact – and our report shares our achievements and lessons learned. Atlassian undertook five pledges through the Pact, which we describe in the report by sharing our key initiatives, achievements, and lessons learned in each area. At a high level, these pledges are: (1) adopting an organizational strategy for AI governance that promotes internal use of AI; (2) addressing high-risk AI in our products and operations; (3) educating our teams about AI and its impacts; (4) providing transparency for deployers utilizing Atlassian AI; and (5) designing our AI systems to ensure that end-users recognize AI in their interactions. While the first three pledges were mandatory for participation in the Pact, the latter two were elective and reflect Atlassian’s responsible technology commitments.

Key lesson: committed teams can drive responsible AI practices across global organizations, but new ways of thinking and doing are required. The underlying theme throughout our report is that committed teams can change organizational practices. From our internal review processes to AI School for Atlassian developers and our product design choices, teams across Atlassian have embraced responsible AI and developed practices that will enable compliance with the [EU AI Act](#). At the same time, organizational evolution requires new ways of thinking and doing. We’ve had to encounter challenges and learn ways to overcome them, only to find a new challenge on the other side. Dialogue between teams thinking about responsible AI strategy and the teams driving responsible AI implementation is critical. Together, we can be the change we seek in responsible use of technology – in Europe and around the world.



Fostering AI adoption across Atlassian through internal governance

Our first pledge under the AI Act was to “[adopt] an AI governance strategy to foster the uptake of AI in the organization and work towards future compliance with the AI Act.” Fortunately, we’ve learned a few things about how to balance responsibility and innovation – and we’re still learning.

Deploying AI across a global team of over 13,000 people is no small feat. Before the first prompt was entered into a Rovo agent, our cross-functional Responsible Technology Working Group – with connections across legal, program management, engineering, and technical disciplines – was hard at work on developing answers to foundational questions about how AI would be used at Atlassian.

Our teammates had questions and we needed guardrails. What are the rules that we need to follow or guidelines we should consider? Just because we can use AI to complete this task, should we use AI given the potential downsides? And how are we going to operationalize these decisions for our fellow Atlassians?

Our approach to AI governance is rooted in our Responsible Technology Principles. In 2023, we published our [Responsible Technology Principles](#). Our Principles are a set of guiding commitments designed to ensure that Atlassian develops and deploys technology – especially AI – in a way that is ethical, transparent, and aligned with our core values. Much of our work on internal governance stems from our Principles and brings them to life through our practices. In many cases, our Principles are derived from our [Atlassian Values](#).

- **Open communication, no bullshit.** Openness is foundational to Atlassian – one of our core values is [Open Company, No Bullshit](#). It's important that anyone who wants to make the most of new technologies is equipped with the right information to do so. In practice, we've applied this principle to drive initiatives like new design features to ensure that it's clear when AI is in-use, as well as transparency disclosures to support AI deployers as they utilize our AI products.
- **Build for trust.** Trust is at the heart of our work and our products: if someone doesn't trust our company, they won't use our products or want to work here. This extends to the technologies that underpin and power our products and our work. We communicate about our trust posture through our [Trust Center](#), which provides in-depth information about our compliance and certification work, as well as our [AI transparency disclosures](#).
- **Accountability is a team sport.** At Atlassian, our mission is focused on collaboration and teamwork. Our products are powered by our own people, upon the foundational technologies that we use to deliver them – and, of course, by how our customers' teams choose to use them. We've applied this principle in our products through features like feedback collectors, which use thumbs-up and thumbs-down indicators, and through external engagement with external AI experts that advise us on responsible AI, like the [Human Tech Institute](#) at the University of Technology Sydney.
- **Empower all humans.** At Atlassian, we want our company and our technologies to be open, inclusive, fair and just: to reflect the human-centric values and fundamental human rights that we all share. Our journey to build responsibly reflects this goal. As we describe in next section, our Responsible Technology Reviews – and our product design choices – help us deliver on this promise.
- **Unleash potential, not inequity.** We know that behind every great human achievement, there is a team. We also believe that new technologies can help empower those teams to achieve even more. If we use these technologies (like AI) responsibly and intentionally, then we can supercharge this vision and contribute to better outcomes across our communities.



Our key actions

Across Atlassian, teams conduct Responsible Technology Reviews (RTRs) to ensure that our AI products and use cases reflect our commitments. RTRs are structured assessments designed to ensure that technologies - especially those involving AI - are developed and deployed consistent with our principles, values, and legal obligations. We have published our RTR template in our No BS Guide to Responsible AI Governance, which outlines the core pillars of our Responsible Technology Program and how we apply our Responsible Technology Principles to AI development and deployment.

Our RTRs have helped us prepare for compliance with the AI Act. We use RTRs across our organization to provide a consistent view of relevant AI use cases and their associated capabilities, limitations, and risks. The same RTR assessment is conducted whether a use case relates to our products, our internal tooling, or our procurement of third party technologies. Because of this, the use of RTRs becomes a core tool for:

- all of our employees working with AI to uplift their AI literacy, and
- teams managing the RTRs to assess use cases for risk, ensure that appropriate safeguards will be applied to limited-risk AI systems, and identify scenarios that could be high-risk without further intervention.



Our achievements

We have conducted over 175 RTRs against our AI products and deployment. We have embedded the RTR process into existing workflows across our organization, including product development, procurement, and internal AI approval processes. The RTR process focuses on ensuring that teams are able to understand, explain, and document the potential impacts, risks, and benefits of their AI projects. It prompts project owners to identify key decisions and trade-offs that they need to consider in order to implement AI use cases responsibly and transparently. Our Responsible Technology Working Group has reviewed each of these submitted RTRs, and works with the submitting team to help them achieve completion and keep RTRs up to date over time.

We have publicly shared our RTR template to inspire peers, support customers, and advise policymakers. When we developed our RTR process, most existing AI impact assessments were too complex for self-serve use. This was especially challenging for Atlassian, where AI use cases span many departments and locations. We created a simpler RTR template for our teams and shared it externally to expand others' understanding of how AI review processes can work, especially in lower-risk contexts.



Lessons learned

Our RTRs provided insight into consistent challenges facing AI developers. Atlassian's AI products and internal use cases typically do not invoke high-risk concerns. However, applying RTRs across the breadth of our AI work surfaced valuable insights about challenges facing AI developers in applying our Responsible Technology commitments. Three primary focal areas emerged:

- **Inclusivity.** Ensuring that our AI technologies are representative of and accessible to a wide audience, including across geographies and languages.
- **No surprises.** Aligning our end user communications and transparency initiatives with the expectations of our audience, whether they are customers or employees.
- **The AI value chain.** Understanding the scope of our accountability in the AI value chain, particularly where we are reliant on upstream providers and their efforts in training, testing, and validating their AI.

Casting a wide net across AI use cases can help inform refinement of governance tools over time. From the start, we asked as many teams to complete RTRs as possible, regardless of the inherent risk of their AI product or use case. This gave us a broad picture of how our teams were using AI across our organization, both in building products and enhancing our operations. With those lessons learned, we adapted the scope of our responsible technology reviews to focus on the most important use cases for each of our teams and departments. This allowed us to maintain a consistent format for all RTRs, while allowing for lighter-touch review and supervision of lower-risk use cases across our organization.

Multifaceted review processes – for privacy, cybersecurity, and other concerns – are key to bolstering AI governance. RTRs are an important tool, but AI governance is interdisciplinary. We rely on a combination of existing processes governing core concerns like privacy and cybersecurity, and RTRs to understand and assess new, AI-specific concerns. This avoids burdening our teams with duplicative processes or questions. For example, we ensure that RTRs do not repeat questions raised in other assessment processes, and that RTRs present open-ended questions designed surface AI-specific risks and impacts.



Addressing high-risk AI in our products and our own AI deployments

Our second pledge under the AI Pact was to “carry out to the extent feasible a mapping of AI systems provided or deployed in areas that would be considered high-risk under the AI Act.”

This framing resonated with our approach to driving adherence to our Responsible Technology commitments, which involves close coordination with teams who are navigating compliance and innovation. Mapping our AI landscape requires human guides with eyes on potential risks.



Our key actions

RTRs are our primary means to ensure that we are not developing – or deploying – AI products for high-risk use cases. We use our RTRs in combination with other assessment processes across privacy, security, and legal to provide a complete picture of the risks and impacts associated with a given AI use case. This allows us to understand whether a use case could be considered high risk, and work with teams to reconsider and rework their projects to avoid this designation.

In addition to RTRs, we also maintain an AI Gateway, which controls access to foundation models for the purposes of prototyping, building, etc. The Gateway applies across both internal and external use cases and requires completion of an RTR for access, ensuring that we’re looking closely at intended use through the RTR – and prior to approval through the AI Gateway process.



Our achievements

We have updated our Acceptable Use Policy to specifically address high-risk AI. As an enterprise software provider, we are clear with our customers that our products cannot and should not be used for certain purposes. In a new AI-specific section of our [Acceptable Use Policy](#) (AUP), we clarified that our AI products should not be used in circumstances that would be considered high-risk under a range of relevant laws or as a result of known AI-specific risks. For example, our AUP prohibits high-impact decision-making, the provision of specialist advice, and activities that would be considered prompt injection.

We have developed our own technical measures to mitigate the risk of high-risk use. AI can help end users avoid high-risk use. When an end user attempts to use our AI products for a purpose that is clearly prohibited by our AUP, our products are designed to return a result that reminds them about our prohibitions. Many of our upstream providers offer similar features, and we have developed our own protections to support our end users in using AI responsibly. We continue to refine our technical approach to these challenges through customer engagement about their needs.



Lessons learned

Educating organizations about high-risk AI is an ongoing challenge. Across our internal and external stakeholders, there are still opportunities for education about what now constitutes high-risk AI and why our products are not suited for those uses. Through new laws like the EU AI Act, the concept of “high-risk AI” has been further developed and codified in law. However, the parameters of high-risk AI are not understood among organizations deploying AI.

For example, when an end user triggers an in-product filter to prevent a prohibited query, they don’t always realize why they’ve been prevented from pursuing that query, or follow the provided link to review the AUP. Some users continue to try variations on the same blocked query, or reach out to customer support to ask why they are unable to proceed with a given prompt. This means that they are learning the parameters of high-risk AI through inefficient (and frustrating!) interactions with AI systems.

3

Raising the bar on AI literacy across Atlassian

Our third pledge under the Pact was focused on AI literacy. Specifically, we committed to “promote awareness and AI literacy of their staff and other persons dealing with AI systems on their behalf, taking into account their technical knowledge, experience, education, and training and the context the AI systems are to be used in, and considering the persons or groups of persons affected by the use of the AI systems.”

For Atlassian, this meant ensuring that we made the right investments in our people to ensure that they were ready to integrate AI into our products and our own operations. Teams took this obligation to heart and developed in-depth internal training programs focused on stakeholders who have hands-on roles in developing and deploying AI at Atlassian.



Our key actions

Our AI School is focused on ensuring that Atlassians develop a deep understanding of AI. Atlassian’s AI School is a comprehensive internal training program designed to upskill employees in artificial intelligence, with a strong emphasis on practical, real-world applications relevant to Atlassian’s products and teams.

- The curriculum is structured into three main tracks—101 (beginner), 201 (intermediate), and 301 (advanced)—to accommodate a wide range of learners, from those new to AI to experienced machine learning engineers.
- Classes cover foundational AI concepts, prompt engineering, building and evaluating AI agents, and advanced topics such as model training, deployment, and optimization. The learning

experience blends self-paced online modules, live tech talks, hands-on workshops, and guest lectures, culminating in a capstone hackathon or Shiplt event where participants apply their learning to real projects.

- Typical participants include engineers, product managers, program managers, designers, and aspiring machine learning engineers who are either working on or planning to work on AI features. The program is global, with sessions scheduled across time zones and local champions supporting regional engagement to ensure equitable access and foster local learning communities.

Our Responsible Technology Training deepens Atlassians' understanding of our commitments. To complement role-based programs like Atlassian's AI School, we developed a targeted training designed to raise awareness among Atlassians of the capabilities, limitations, and risks associated with AI. Our Responsible Technology Training is designed for any Atlassian who is using AI in their day-to-day work or interested in how AI can be built, deployed, and used more responsibly. This training module is designed to help employees understand and spot common AI risks and issues as they may arise in their work, regardless of their role.



Our achievements

AI School is extremely popular among Atlassians. AI School has grown year-over-year, reflecting strong demand for AI education. Roughly 90% of participants report that AI School advanced their understanding of AI and that they would recommend the course to fellow Atlassians. The program's success is measured not only by these statistics but also by the increasing number of Atlassians who contribute to AI-driven innovation across the company.

Collaboration and experimentation is essential to reaching wider audiences. Our Responsible Technology Training was initially made available internally as a rough pilot program using our own Confluence and Loom products, rather than being made available on our standard learning platform. This allowed us to monitor views, understand engagement metrics, and respond directly to feedback before formally launching the training. Given the intended broad reach of this training, the opportunity to iterate and improve was invaluable.



Lessons learned

Our internal programs provided insight into challenges in AI adoption.

Atlassian's AI School program has been a key part of our broader AI rollout, aiming to upskill employees and foster a culture of AI-driven innovation. Through our engagement, we've learned more about some of the challenges that broad adoption of AI can present, even among people who work in technology.

- **Resistance to change.** Some employees were hesitant to adopt new AI tools and workflows, fearing disruption to established routines or concerns about job security. Champion networks and leadership advocacy helped overcome these concerns.
- **Integration complexity.** Integrating AI into existing products and workflows proved challenging, especially given the need for deep product integration and cross-tool capabilities. This required technical investment and cross-team coordination.
- **Training and adoption.** There were requests for more tailored training to meet the diverse needs of different roles and departments. Ensuring high adoption rates meant providing a mix of hands on workshops, on-demand learning, and targeted sessions. As a distributed organization, we have also leveraged asynchronous tools and ways to share knowledge, including role- and department-specific Slack channels.
- **Customization needs.** Departments needed AI solutions tailored to their specific workflows, which required flexible tools and support for custom AI agent integration.



Supporting our customers in their AI deployments

AI deployers play a critical role in the AI value chain by bringing models and systems to end users. Our fourth pledge under the Pact is to “inform deployers about how to appropriately use relevant AI systems, their capabilities, limitations, and potential risks.” This commitment dovetailed with our ongoing investments in supporting our customers with their deployments, including transparency disclosures about products.



Our key actions

Our Atlassian University program supports AI deployers using our AI products. Atlassian University offers a growing suite of AI-focused courses designed to help users and teams harness the power of artificial intelligence within Atlassian’s cloud products. These courses are accessible online and are structured to support a wide range of learners, from those new to AI to experienced professionals seeking to deepen their expertise.

Courses are delivered as self-paced online modules, often supplemented by hands-on labs, quizzes, and opportunities to earn digital badges or certificates upon completion.

Typical participants include IT administrators, project managers, software engineers, business analysts, and team leads—essentially anyone responsible for configuring, managing, or using Atlassian products in their organization. Many attendees are professionals seeking to upskill in AI to drive digital transformation and efficiency within their teams.

Our transparency disclosures provide key technical information to deployers. [Atlassian's transparency disclosures](#) are structured documents designed to provide clear, accessible information about how Atlassian's generally available AI-powered features work.

Transparency disclosures help users, customers, and regulators understand the technical details, limitations, and intended use cases of Atlassian's AI features. A typical transparency disclosure includes:

- How the feature uses AI models, including model providers
- A description of the AI feature and its intended purpose
- Limitations and considerations on how best to use the feature
- How the feature uses and handles customer data, including data retention and permissions
- Links to further resources and support documentation

We tailor information to our customers and user audiences. In addition to the more detailed technical and learning documentation, it is important that our customers and users have the tools and information they need to quickly adopt our AI-powered features.

We have published plain language adoption guides targeted at two core user types: [administrators](#) and [end users](#).

These guides are further bolstered by the wide range of learning sessions, webinars, and breakout discussions made available to our users. These sessions are led by subject matter experts within Atlassian and hosted at our twice-annual flagship events, [Team](#) and [Team Europe](#), via our [Atlassian Community Events](#) in cities around the world and on-demand online.



Our achievements

We have publicly shared our learning resources to help develop good AI adoption practices. At Atlassian, we have a long track record of sharing the types of practices that help teams work better together, in support of our company mission. We call these Plays, shared in our [Team Playbook](#). We are continuing to share AI-specific learning plays to help our customers and others learn about the capabilities and limitations of AI, such as:

- [AI Training Workshop](#): How to run an interactive training workshop to help team members upskill by learning basic, relevant ways to use AI for their specific jobs.
- [AI Teammate](#): How to help teammates build their first AI agent, specifically tailored to meet team members' unique needs and workflows.
- [AI Innovation Day](#): How to run a dedicated day for learning and practicing AI to enhance teams' skills. Teams can improve their understanding of AI by learning how it operates, identifying practical use cases, and experimenting with various AI tools.



Lessons learned

Discoverability of resources is as important as accessibility of content. Our initial approach focused on ensuring that our resources for deployers were as user-friendly as possible, presented in plain language and tailored to the different roles and experiences within our user base. But across our customer base, we found examples of customers not being able to easily locate this information when needed. We are increasingly focused on curating and enhancing the discoverability of our resources.



Ensuring awareness of when AI is utilized

Our principled commitment to empowering all humans with technology necessitates that people know when they're using AI, especially as AI utilization increases. Consistent with our Responsible Technology Principles, our fifth pledge under the Pact is to “design AI systems intended to directly interact with individuals so that those are informed, as appropriate, that they are interacting with an AI system.” We’ve operationalized this commitment through our products and the interfaces that end-users interact with everyday.



Our key actions

Our AI-powered products and features share a common visual language. We have established and continue to evolve our user interfaces to help end-users clearly identify where we use AI systems within our products. These design choices collectively highlight and allow for differentiation of our AI-powered features. We are also careful in how we apply these elements across the many surfaces where AI systems are used to avoid overwhelming users. The current design elements include:

- **Logos:** Using consistent logos for our AI-powered features and Rovo product whenever they appear in our product suite.
- **Icons:** Employing consistent but visually differentiated icons for AI-powered features, for example featuring star shapes within the icons used for our Rovo Chat and other AI features.
- **Border:** Displaying a rainbow-colored border in scenarios where AI is in use inside the product.

- **The phrase “AI”:** Explicitly identifying features as AI when invoked by a user, for example by typing the shortcut /ai to invoke the Rovo dialogue box or selecting “Ask AI” from a floating toolbar.
- **Footer:** Displaying the footer “Uses AI. Verify results.” persistently in the Rovo Chat window, or after Rovo returns responses in-product.



Our achievements

Experimentation led to a punchier, more proactive disclosure. Our content design team has revisited and revised the footer used in our AI-enabled products to notify users that AI is in use over time. Through experimentation and internal collaboration, we moved from a more passive and less specific disclosure (“Powered by AI. Content quality may vary.”) to a shorter, more active disclosure (“Uses AI. Verify results.”). This newer disclosure can be used across a variety of surfaces where space is at a premium, while also being clearer to users about how best to engage with the potential risks associated with the underlying large language models, like hallucinations.



Lessons learned

We need to stay ahead of future notice fatigue. Our current footer to notify users when AI is in use is workable because it is short and action-oriented. But as AI systems proliferate across enterprises and the EU AI Act’s transparency obligations come into effect, we know that users will increasingly dismiss repetitive and unhelpful notices and warnings, however valuable.

Customer research and insights are essential to success. Atlassian has embraced AI, both within its products and internally. As a result, it is tempting to believe that AI has by now become as familiar to our customers as it is to our employees – and that it is no longer necessary to explicitly identify specific AI-powered features or products. Through our research, insights, and customer relationships, we are ensuring that we are still able to support and identify with customers who are in an early stage of their AI adoption journeys.

Concluding thoughts

Atlassian is proud to support the EU AI Pact. We appreciate the European Commission's initiative to convene the Pact and invite participation from across the AI value chain, including companies like Atlassian that provide a critical layer of innovation between developers of foundation models and deployers of AI systems. We value opportunities to provide insights into our Responsible Technology initiative with our stakeholders, and we look forward to sharing more about our preparations for the EU AI Act as its implementation proceeds.

Interested in learning how Atlassian develops Rovo AI responsibly?

Review our [No BS Guide to Responsible AI Governance](#)

