



FULL POLICY BRIEF

Ethical and Legal Impact Assessments and Citizen Participation in Technology: Final Recommendations

Deliverable D.6.4 | Horizon Europe | September 2025

Abstract

For three years, VUB conducted annual **Ethical and Legal Impact Assessments of technologies (ELIAs)** for the TITAN project. These assessments are part of a complex regulatory landscape that **assess the risks of technologies**. In parallel, the project engaged regularly with citizens in co-creation processes to ensure user requirements were met.

The TITAN project – which **develops an AI conversational agent designed to combat online disinformation**– provided an ideal place to explore how these assessments are implemented in practice.

This policy brief offers lessons learned from the TITAN project directed to **industry and policymakers**, highlighting ethical and legal recommendations of ELIAs, best practices, citizen involvement, and governance challenges.

From this practice, it identifies critical gaps in existing legal and ethical frameworks, particularly around standardization and interdisciplinary collaboration. Findings emphasize that **impact assessments of technologies are continuous processes requiring translation across technical, legal and societal domains**.

This brief provides critical steps on how to integrate ELIAs into AI development, ensuring responsible innovation while safeguarding democratic values and fundamental rights.

Key Points

- **Make ELIAs practical:** translate complex ethical and legal assessments into clear actionable steps for the technology development.
- **Standardized methodologies:** develop harmonized guidelines for conducting ELIAs in AI projects.
- **Clarify purpose:** define how Legal Impact Assessments (LIAs) should be effectively conducted under the AI Act and related frameworks.
- **Ensure continuous relevance:** treat impact assessments as central and ongoing, adaptive process that evolve with AI systems.
- **Bridge disciplines:** strengthen collaboration between technical, legal, ethical and social disciplines through facilitated dialogue.
- **Engage citizens:** involve end-users early in the development process.

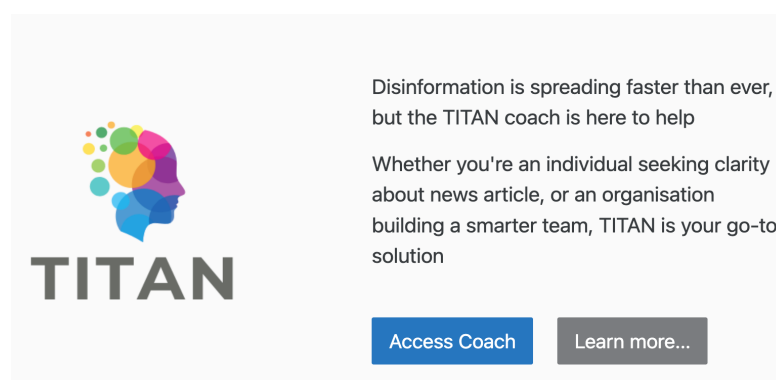
Introduction to TITAN: An AI Conversational Agent to Fight Disinformation

Disinformation is accelerated by Artificial Intelligence (AI), destabilizing democratic societies, eroding trust in institutions and amplifying false narratives online. In a digital landscape where opinions often spread faster than facts, both individuals and societies need tools to identify misleading content before it spreads.

AI presents both challenges and opportunities in this fight. While AI can amplify disinformation at a large scale, it can also detect misleading content, counter false narratives, and engage users in ways that foster critical thinking to prevent disinformation. Scientific literature¹ suggests that the use of AI in content moderation should not be done without also prioritizing human review processes. In other words, to combat disinformation, the focus should also be placed on people, by giving them skills to recognize disinformation *before* they spread it.

In response to this challenge, the **TITAN project developed an ethical AI Conversational Agent designed to combat disinformation**. TITAN integrates multiple systems to provide effective interactive engagement:

- **AI-driven dialogues** built on pre-established ML techniques, as well as LLMs, allowing users to actively exchange views on written comment, including news or social media posts with the conversational agent.
- **Disinformation signal detection** flags and ranks content to guide TITAN's conversational agent's responses, ensuring that discussions focus on potentially misleading or false information.
- **Socratic questioning methodology** to encourage reflective thinking. The conversational agent was trained with Socratic questions; by incorporating this methodology the system improves users' critical thinking by coaching them.
- **Critical Thinking Assessment tests** that evaluate users' responses and understanding, providing feedback on user's current critical thinking capabilities. This feedback informs subsequent dialogues and allows the system to adapt to better support individual learning.



¹ Marsden, C., Meyer, T., & Brown, I. (2020). Platform values and democratic elections: How can the law regulate digital disinformation?. *Computer law & security review*, 36, 105373.

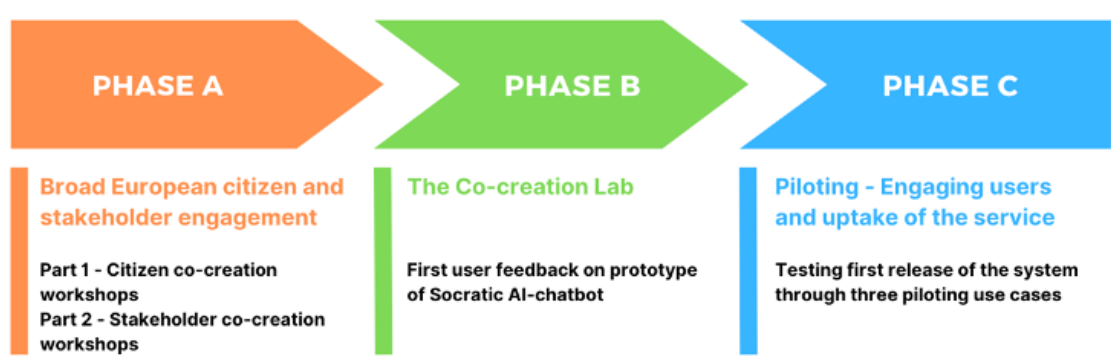
Through this approach, TITAN not only counters disinformation but also fosters a better informed and reflective digital society. The TITAN project exemplifies how AI plays a role in the fight against disinformation, while also demonstrating the need to consider the broader socio-political context in which such AI technologies are developed and deployed. Their impact extends beyond technical design, raising questions of citizens manipulation, technological influence, and democratic resilience.

TITAN technology has been assessed through Ethical and Legal Impact Assessments (ELIAs) to ensure compliance with fundamental rights and ethical principles. Striking this balance between innovation and citizens' rights is essential for strengthening a democratic digital society.

The TITAN project combined two complementary approaches conducted in parallel. One is the **co-creation** process with users, involving citizens and civil society actors in the design and testing of the conversational agent. At the same time, we conducted annually an **Ethical and Legal Impact Assessment (ELIAs)** to navigate the complex and evolving regulatory landscape, identifying risks and ensuring compliance with ethical and legal standards. The ELIAs were also supervised by an External Ethics Board of experts. Together, these steps bridge regulatory requirements, ethical assessments, and citizen's opinions.

Socio-Political reflections and Recommendations: The co-creation Journey

The co-creation process within the TITAN project brought together researchers, developers, civil society, and end-users to define priorities for the AI chatbot. Across three phases, the Socratic AI system was co-created and validated with **citizens and stakeholders** in multiple European countries. Phase A engaged citizens and experts to identify and validate initial user requirements, Phase B refined these requirements through Living Lab prototype testing with diverse groups (students, NGO members, citizens, and migrants), and Phase C piloted the system in higher education, NGOs, and migrant communities, generating final insights on usability, functionality, and its impact on critical thinking. This collaborative approach ensured that technical decisions were also informed by users' needs, social realities, ethical considerations, and legal obligations.



From the early stages, the project ensured that user requirements were consistently translated into tangible outcomes and functionalities. Citizen and stakeholder feedback

highlighted the need for clarity and transparency in how the system operates, protection of personal data in accordance with the General Data Protection Regulation, avoidance of ideological bias, and accessibility for users with diverse levels of skills and knowledge. These requirements were implemented and assessed through our three iterative testing phases, highlighted above. Plain-language explanations of system functions were provided, supported by multiple rounds of testing to refine the information presented, offer language options, and ensure the system was clearly understood. The co-creation framework allowed these principles to be revisited at each development stage, ensuring alignment between user expectations and technical delivery.

The process also revealed critical socio-political issues that extend beyond the technical sphere. **Public trust** emerged as a central concern, dependent on both system reliability and its perceived independence from political or commercial influence. The potential **ideological impact** of the tool was recognised as a double-edged risk: while it can strengthen critical thinking, it may inadvertently amplify certain narratives. **Informed participation** was identified as essential, requiring that all users understand the system’s purpose, benefits, and limitations. Transparency is therefore important. Finally, questions of **responsibility and accountability** were raised, underlining the need for clear mechanisms to address possible harms arising from AI recommendations. For example, if the system were to misclassify information, inadvertently reinforce misleading narratives, or contribute to a decline in trust in reliable institutions.

To address these findings, the following stakeholder-oriented actions were recommended:

Stakeholder Group	Recommended Action
Policy-makers	Establish mandatory transparency standards for AI systems that influence public discourse, including disclosure of bias mitigation strategies and data sources.
Developers	Integrate multidisciplinary ethical reviews throughout the development cycle and publish accessible impact assessments.
Civil society	Monitor ideological neutrality and advocate for the inclusion of marginalised groups in both design and testing phases.
End-user organisation	Provide training and guidance to ensure critical engagement with online AI outputs, particularly in vulnerable or digitally excluded communities.

The TITAN co-creation journey demonstrates that ethical and legal compliance is only the foundation of responsible innovation. **Long-term trustworthiness depends on sustained engagement with diverse stakeholders, proactive mitigation of ideological risks, and a commitment to transparency and accountability at every stage of the development process.**

Ethical and Legal Recommendations: A Necessary Step Forward

Given their complexity and novelty, AI systems cannot be responsibly deployed without first addressing compliance deficiencies and regulatory uncertainties. Our work highlights the

urgent need for **clearer legislation in practical ethics and legal risk assessment guidelines of technology** to support effective implementation.

Existing impact assessments are often complex and fragmented, spread across multiple ethical approaches and legal assessment frameworks. A necessary first step is to map which one is relevant for the deployed AI. To support this, we reviewed the most relevant European Union standards for ethical and legal impact assessment of technologies (ELIAs). Below, we present a summary of our mapping:

Summary of existing Ethical and Legal Impact Assessments of Technologies (ELIAs)

Ethical Assessments of Technologies – Guidelines (non-binding)

- AI Principles ² – OECD
- Guidelines for an Ethical Use of AI & AI Impact Assessment Tool³ – SHERPA project
- AI Ethics Guidelines: European and global Perspectives⁴ – CAHAI
- Ethics Guidelines for Trustworthy AI⁵ – HLEG on AI
- AI and Robotics: Ethical Framework⁶ – SIENNA project
- Ethics by Design and Ethics Use Approaches for AI⁷ – EU Commission DG Research and Innovation
- General-Purpose AI Code⁸
- Recommendations on the Ethics of AI⁹ – UNESCO
- Methodology for the risk and impact Assessments of AI Systems¹⁰ – CAI

Legal Impact Assessment of Technologies – Binding when required

- Data protection Impact Assessment (DPIA)¹¹ – (Art. 35) GDPR
- AI Fundamental Rights Impact Assessment¹² – (Art. 27) AI ACT
- Online Platform Risk Assessment^{13 14} – (Art. 34) Digital Service Act (DSA)
- Gatekeeper Compliant Assessment¹⁵ – (Art. 7) Digital Markets Act (DMA)

² <https://oecd.ai/en/ai-principles>

³ <https://project-sherpa.eu/wp-content/uploads/2019/12/development-final.pdf>

⁴ <https://rm.coe.int/cahai-2020-07-fin-en-report-ienca-vayena/16809eccac>

⁵ <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

⁶ <https://www.sienna-project.eu/w/si/404>

⁷ https://ec.europa.eu/info/funding-tenders/opportunities/docs/2021-2027/horizon/guidance/ethics-by-design-and-ethics-of-use-approaches-for-artificial-intelligence_he_en.pdf

⁸ <https://digital-strategy.ec.europa.eu/en/policies/contents-code-gpai>

⁹ <https://unesdoc.unesco.org/ark:/48223/pf0000381137>

¹⁰ <https://rm.coe.int/cai-2024-16rev2-methodology-for-the-risk-and-impact-assessment-of-arti/1680b2a09f>

¹¹ <https://ec.europa.eu/newsroom/article29/items/611236>

¹² <https://artificialintelligenceact.eu/article/27/>

¹³ <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32022R2065&qid=1666857835014>

¹⁴ https://www.eu-digital-services-act.com/Digital_Services_Act_Article_34.html

¹⁵ https://commission.europa.eu/strategy-and-policy/priorities-2019-2024/europe-fit-digital-age/digital-markets-act-ensuring-fair-and-open-digital-markets_en

The analysis of ELIAs revealed overlaps but also significant gaps and inconsistencies, particularly in how risks are defined, assessed and communicated to different stakeholders. Risk categories are often too general and considered subjective in conversations with engineers such as “fairness of AI” or “human agency and oversight”. ELIAs are **perceived by technicians as burdensome and complex process, limiting their ability to generate clear, actionable steps for implementation.**

Since the central question is: **how can ethical and legal impact assessments be conducted effectively to guide responsible AI deployment?** our review confirms that in practice, while these assessments are necessary, they continue to face recurring obstacles that must be addressed.

First, the TITAN project offered concrete insights into the specific challenges of assessing an AI conversational agent designed combat misinformation:

Project-Specific Guidelines		
Issue	Description	Recommendations and links to general guidelines
System differentiation	TITAN’s AI comprised multiple subsystems with distinct risks: the conversational agent raised risks of e.g., user manipulation; the critical thinking test presented accuracy risks and transparency; dialogue creation involved data-gathering risks. Current ELIAs rarely account for such subsystem differences.	ELIAs should explicitly differentiate between subsystems to capture distinct technical, ethical, and legal risks.
Risk of unintended manipulation	Existing guidelines did not adequately address unintended manipulation, a risk identified by TITAN’s external ethics board in relation to conversational agents.	Expand ELIAs to explicitly consider unintended manipulation risks, especially for interactive AI systems.
Interdisciplinary gaps	During our conversations, technical experts struggled to grasp the purpose, process, and results of ELIAs. They also needed explanations of basic legal concepts (e.g., “personal data” under GDPR). As well as for them to translate complex AI and technical systems into coherent summaries to be used for the assessments. Many risks were “to be determined” at early development stages, which added uncertainty.	Provide simplified ELIAs, supported by facilitators, and foster stronger collaboration between technical, legal, and ethical experts throughout the development cycle.
Citizen’s involvement	Citizen workshops showed alignment with ELIAs on privacy/data protection but also revealed concerns on trust, usability, and inclusivity (e.g., marginalized groups, language barriers). These issues were not captured by ELIAs alone.	Integrate participatory design and user feedback into both technology development and ELIAs to ensure societal concerns are reflected.

Complexity and length	TITAN's final ELIAs spanned 61 pages, which required extensive summarization into actionable outputs for technical partners.	Streamline ELIAs into structured, accessible outputs tailored for technical implementation while keeping full assessments available for oversight and dialog.
Iterative assessment process	TITAN performed annual ELIAs, supervised by an External Ethics Board. The project combined a Data Protection Impact Assessment (DPIA) with the HLEG Guidelines on Trustworthy AI.	Combined legal and ethical assessments, supervised by external experts and offer facilitators roles, to adapt to evolving risks and regulations. Facilitate dialog between different experts.

Second, the three-year deployment and analysis of the assessment allowed us to derive a set of general recommendations for future use:

General Guidelines		
Issue	Description	Recommendations
Lack of standardization	Current ELIAs remain largely undeveloped or inconsistent; even existing tools (e.g., DPIA, AI Fundamental Rights IA) are complex and not fully operationalized. Creating uncertainty in how AI systems should be assessed.	Develop standardized, practical ELIA methodologies with clear, actionable guidance for practitioners.
Complexity and Length	Conducting ELIAs is lengthy and resource intensive. Requiring technical experts to simplify complex AI systems, while ethical/legal teams often fail to translate values into practical steps.	Simplify outcomes into structured, implementable steps; provide clear communication tools and dialogs for both technical and non-technical teams.
System Differentiation	AI systems often consist of subsystems, that can be developed separately, with distinct risks (e.g., data use, manipulation, accuracy). Current guidelines rarely account for these differences, and attempt to assess AI as one.	Require subsystem-level in AI assessment within ELIAs, to ensure risks and specific ethical concerns are captured in detail.
Subjectivity of Risks	Ethical/legal risks are interpreted differently by stakeholders, creating variability and ambiguity in assessments.	Establish structured dialogue and offer definitions to align understanding across disciplines.
Interdisciplinary Gaps	Collaboration between legal, ethical, and technical experts remains limited; many technical experts lack ELIAs knowledge on its purpose and	There is a clear need for explainability of both technical systems and ELIAs purposes

	usefulness, as well as legal and ethical experts do not comprehend complex technical systems.	and values to ensure meaningful assessments. Ensure dialog and explainability of both technical systems and ELIAs; promote closer interdisciplinary collaboration.
Communication Barriers and The Role of Facilitators	The quality of the output depends on the quality of the input. Without facilitation, technical and legal experts struggle to bridge their knowledge gaps. They are crucial both for explaining the technology and translating ethical and legal findings into actionable steps.	Integrate facilitators to mediate discussions, explain technical systems, and translate legal/ethical findings into actionable steps. Promote structured dialogue between engineers, legal/ethical experts, and social scientists.
Lack of citizen participation	Users' concerns, risks, and fears are often absent from current ELIAs.	Make participatory design an element of ELIAs, not only for system design but also for identifying and mitigating risks.
Limited Ethical Focus	Non-binding ethical assessments are often overlooked despite their added value.	Encourage inclusion of ethical evaluations (e.g., HLEG guidelines on Trustworthy AI) alongside legal ones to guide a more sound and responsible AI development.
Process, not checklists	Effective assessments involve ongoing dialogue between technical experts and facilitators, allowing risks to be identified, balanced and mitigated throughout the AI development lifecycle. ELIAs are often treated as checklists rather than iterative processes.	Adopt ELIAs as ongoing assessments, integrating continuous dialogue and balancing risks across AI lifecycle.

Thirdly, based on TITAN's Horizon project creation of an AI conversational agent, and the ethical and legal impact assessments used to assess the technology, we recommend:

Policy Recommendations Acting on Ethical and Legal Insights from TITAN
• Standardize ethical and legal frameworks: harmonize EU and international standards and guidelines for AI to reduce ambiguity and ensure consistent assessment of risks.

- **Embed facilitation roles in AI projects:** appoint facilitators to translate complex ethical and legal requirements, and technical explainability, enabling actionable and effective implementation.
- **Offer courses:** these can bridge the disciplinary gap, in which experts can learn about other fields between technical, legal, and ethical expertise.
- **Engage citizens in co-creation:** include citizens in technical development to align AI design with societal values, and to ensure legitimacy and trust in AI governance.
- **Treat impact assessments as living processes:** ensure Ethical and Legal Impact Assessments guide continuous action, through practical steps and follow-up actions, to be embedded into project design, rather than being a one-off compliance exercise.

Lessons Learned from Integrated best practices

Translating insights from citizens' co-creation into actionable steps for technology development required not only collecting citizen input but also embedding it meaningfully into design decisions, ensuring that the final system reflected real user needs while meeting ethical and legal standards. The process demonstrated that effective translation of requirements relies on continuous dialogue between technical experts, social scientists, and users, supported by structured feedback loops and transparent decision-making, which is a time-consuming process.

ELIAs operate in an interdisciplinary environment, which proved both challenging and highly valuable. Differences in terminology, priorities, and working methods initially slowed progress, but ultimately enriched the process by bringing together different expertise. The social science perspective ensured that socio-political concerns, such as inclusivity, transparency or trust were considered alongside technical performance. In turn, technical expertise helped assess the feasibility of proposed solutions and adapt them within existing technical reality.

Adopting a broad perspective throughout the project was essential. This meant not only addressing immediate technical goals but also reflecting on the societal implications of the technology. One of the key challenges was to integrate ELIAs feedback into different development stages. This required careful facilitation to balance diverse viewpoints and priorities, manage expectations, and translate qualitative feedback into concrete design requirements. Finally, integrating citizen input was at times complex, it proved invaluable for identifying blind spots, improving usability, and enhancing public trust in the technology.

A Call to Action

The TITAN project began before the emergence of large-scale generative AI such as ChatGPT, yet AI is now transforming societies at a pace that is difficult to predict. While assessments like DPIA (GDPR) have set important foundations, it falls short in addressing the complex, evolving challenges of novel AI systems embedded in nearly every aspect of citizen's lives.

New regulations also demonstrated how risks are open to change – with systems potentially shifting from low to high risk as technology and its applications evolving.

Policy makers must develop clear, standardized methodologies for Ethical and Legal Impact Assessments (ELIAs) and include citizen involvement in AI design. At the same time, industry should embrace ELIAs as part of innovation, working with facilitators to bridge technical, legal and ethical domains, to build more responsible, fairer, and just technologies.

Ethical and Legal Impact Assessments (ELIAs) are essential tools for AI governance, but they remain in an early stage of development. The TITAN project demonstrates both the value of these assessments and the significant gaps that persist in current frameworks. Without clearer standards and stronger interdisciplinary collaboration, risk assessment is becoming either too abstract or too burdensome to guide real-world innovation and evaluation. Yet, Ethical and Legal Assessments refer to different values, coming from two different disciplines. Our role in the TITAN project was to apply both to evaluate a complete risk assessment, that goes beyond what is legally required. Our evaluation highlights the lack of clarity and standards to evaluate complex AI systems.

Our main message is this: **AI will shape the future of democracy. Ensuring that it does so responsibly requires dialogue among different stakeholders. Complexity should not become a barrier for developers in applying ethical and legal standards or assessing their AI models.**

Future regulation must be adaptive and realistic to technological change, by keeping pace of complex technological development while safeguarding citizens' rights. Lawful AI is not only about compliance with rules – it is about responsibility. To achieve this, ethics guidelines must also play a role, and must be translated into actionable steps, supported by supervisory boards that can oversee implementation, but also technical experts who can explain AI's complexities and citizens who ensure alignment of the AI with societal needs and values, and facilitators who can bridge the gap between technicians, lawyers, ethical experts, and society. Only through collective effort AI can be developed in ways that strengthen, rather than undermine, democracy.

Authors

Ana Fernández Inguanzo, VUB, BSOG, CD2I
Aline Duelen, VUB, SMIT
(dr.) Trisha Meyer, VUB, BSOG, CD2I
(dr.) Wendy van den Broeck, VUB, SMIT

Acknowledgements

This document is a deliverable of the TITAN project, which has received funding from the European Union's Horizon 2020 Programme under Grant Agreement (GA) #101070658 and by UK Research and innovation under the UK governments Horizon funding guarantee grant numbers 10040483 and 10055990.