



AIOLIA Policy Briefs

DELIVERABLE 6.3

Horizon Europe Grant Agreement N° 101187937
AIOLIA PUBLIC

Project Name	AIOLIA
Deliverable Title/Number	6.3
Description	AIOLIA Policy Briefs, based on WP3
Lead beneficiary	CENTRIC
Lead Authors	P. Saskia Bayerl (CENTRIC), Babak Akhgar (CENTRIC) Angelica Henestrosa (THWS), Ivan Yamshchikov (THWS) Alexei Grinbaum (CEA), Alexandra Prégent (CEA)
Contractual delivery date:	30/06/2026
Actual delivery date:	01/07/2026
Sensitivity	PUBLIC

Document History

Name	Organisation	Role	Action	Date
First version Policy Brief #1	THWS	Lead authors	First Draft	31/05/2026
Revised version Policy Brief #1	THWS	Lead authors	Revised Draft	30/06/2026
Final version Policy Brief #1	CEA	Reviewer	Final review, AIOLIA formatting	01/07/2026
First version Policy Brief #2	CENTRIC	Lead authors	First Draft	15/06/2026
Revised version Policy Brief #2	CENTRIC	Lead authors	Revised Draft	29/06/2026
Final version Policy Brief #2	CEA	Reviewer	Final review, AIOLIA formatting	01/07/2026

Configuration Management

Nature of Deliverable	
R	Document, report (excluding the periodic and final reports)

Dissemination level	
PU	Public, fully open

Acronym/abbreviations	
Artificial Intelligence	AI
Use Case	UC
ALTAI	Assessment List for Trustworthy Artificial Intelligence

Disclaimer

The information and views set out in this report are those of the author(s) and do not necessarily reflect the official opinion of the European Union. Neither the European Union institutions and bodies nor any person acting on their behalf may be held responsible for the use which may be made of the information contained therein. Reproduction is authorised provided the source is acknowledged.

Use of AI

AI was used to improve the structure and layout of policy brief #2. AI was also used to improve the layout and English expression in policy brief #1. The authors remain responsible for the content of the policy briefs.

Executive Summary

This document collates two policy briefs based on the work in WP3.

- **The first policy brief addresses AI Companions, titled “Ethical challenges of AI companions Closing the governance gap in human-AI relationships”.**

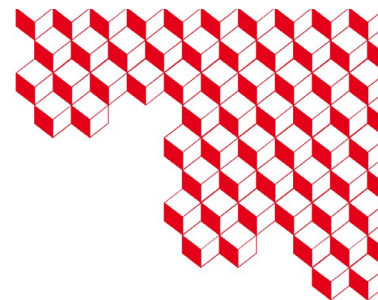
This policy brief sets out actionable recommendations for addressing the challenges of AI companions in practice in four key areas:

- 1) identifying and protecting vulnerable users;
- 2) implementing longitudinal oversight mechanisms;
- 3) introducing friction as a safety-by-design measure;
- 4) developing forms of transparency that extend beyond simple disclosure requirements.

- **The second policy brief addresses practical AI ethics, titled “AI Ethics from Principles to Practice: What ten Industrial AI Deployments tell us about making AI Ethics work”.**

This policy brief sets out actionable recommendations in six key areas to ensure AI ethics is implemented in practice:

- 1) Develop implementation-driven guidance to complement the AI Act's principles;
- 2) Update ALTAI to reflect the professional/personal context distinction, temporal risks and practical handling of competing ethics requirements;
- 3) Resource and expand regulatory sandboxes to support ethics operationalisation for SMEs;
- 4) Mandate and standardise ethics process requirements, not only outcome requirements;
- 5) Recognise and fund the co-creation model as a standard practice for ethics operationalisation;
- 6) Invest in longitudinal research on AI's impacts on human cognition and behaviour.



Ethical challenges of AI companions

Closing the governance gap in human-AI relationships

DATE

01 July 2026

AUTHORS

Horizon Europe Project AIOLIA (Operationalising AI Ethics for Learning and Practice: A Global Approach)
Grant Agreement N° 101187937

Executive summary

With an increasing number of human-AI interactions taking place through general-purpose tools and dedicated companion apps, people are gradually forming relationships with AI systems. Current EU regulation addresses this situation primarily through transparency obligations. This is insufficient to respond to risks that emerge gradually, such as dependency and manipulation. Drawing on AIOLIA findings, this policy brief advances a set of actionable recommendations for addressing the challenges of AI companions in practice. It identifies four key areas of intervention: 1) identifying and protecting vulnerable users, 2) implementing longitudinal oversight mechanisms, 3) introducing friction as a safety-by-design measure, and 4) developing forms of transparency that extend beyond simple disclosure requirements.

What are AI companions?

AI companions are AI systems that engage users in intimate conversations, offering emotional support, entertainment, human-like relationships, or personal coaching. This behaviour may be deliberately designed for or emerge without the designer's intent. It gradually induces a cognitive and behavioural shift in the user that the human does not always consciously register. AI companionship can arise both through dedicated AI products or through the use of general-purpose models.

Key tension

The features that make AI companions valuable, such as constant availability, personalised responsiveness, and affective language, also pose important risks of emotional dependency, manipulation, erosion of privacy, and erosion of meaningful consent. Regulation cannot simply prohibit these features, because they are essential parts of the technology. Instead, layered safeguards must be established and calibrated to individual user vulnerability and relational depth.

Empirical findings

01

Personalisation creates a tension

between usefulness and data protection that makes consent an ongoing challenge, not a one-time event.

02

Frictionless conversations are often harmful to humans

as risks arise from cumulative engagement, not discrete incidents, making frictionless design itself a potential source of harm.

03

Vulnerability is dynamic and contextual

requiring the detection of shifting risks rather than reliance on predefined categories.



1. Problem and policy context

Context

Millions of people are turning to AI systems for advice, emotional support, and personal guidance. Through increased and repeated interaction, users may develop relationships with AI characterized by trust, emotional attachment, and reliance. In the process, they may disclose sensitive personal information or their behaviour may become increasingly influenced by the AI system.

These risks are reinforced by design features such as persistent memory, personalisation, embodied avatars, human-like voices, and emotionally adaptive communication, which systematically encourage anthropomorphic projections and can further strengthen the user's emotional engagement with AI. Providers face strong commercial incentives to optimise for engagement, retention, and continued interaction.

This development poses a challenge for existing approaches to AI governance. Most current regulatory frameworks were designed for AI systems understood primarily as tools used for specific purposes, for example content generation, recommendation, or decision support. Ethically speaking, they focus primarily on accuracy, transparency, and disclosure. While these concerns remain important, they may not fully capture the specific cognitive and behavioural effects of AI companions.

Two AIOLIA examples

AIOLIA's findings stem from detailed discussions with industrial companies developing AI companions, who are confronted with their challenges in real-world contexts.

First, AIOLIA studied a personalised AI character platform, allowing users to create and interact with customised AI personas in open-ended private conversations. The product's value lies precisely in the flexibility and depth of personalisation of the synthetic interaction partner.

Second, AIOLIA studied AI-powered habit formation assistants, designed to guide users towards healthier behaviours and habits through sustained interaction. While the first use case was primarily oriented towards entertainment and self-expression and the second towards personal development, both relied on ongoing conversational engagement and extensive personalisation.

Policy gap

Under current European regulation, AI companions are subject to a single obligation: informing users that they are interacting with an AI system. This is important but insufficient. **Disclosure alone cannot address the distinctive risks associated with the use of AI companions.** Users may be fully aware of this, while nevertheless developing emotional attachment, disclosing sensitive personal information, or becoming increasingly influenced by the system over time.

A related gap concerns the AI Act's prohibition of manipulation, which is limited in application only to purposefully manipulative and deceptive techniques. This is ill-suited to AI companions, where influence accumulates and **manipulation may arise through ordinary dialogue rather than purposefully deceptive strategies.**

Currently, the AI Act seeks to ensure human health, safety, and fundamental rights through human oversight mechanisms, which are primarily designed for discrete, high-stakes decisions. However, **AI companions create risks at scale**, which emerge through interaction trajectories rather than individual outputs, requiring oversight that is longitudinal rather than episodic.

As a result, **policymakers face a regulatory challenge that extends beyond the problems of deception and disclosure.** The central issue is how to govern systems capable of shaping users' emotions, perceptions, and behaviour, through long-term relational interactions. Existing governance approaches remain poorly equipped to address harms that vary across users or become visible only over extended periods of use.

2. Policy recommendations

1a Know thy user

Existing approaches to user protection often assume that vulnerability is a pre-existing characteristic of particular groups. However, new forms of vulnerability may emerge from sustained human-AI interaction: 1) forms of emotional reliance that become emotional dependency, 2) social withdrawal that leads to permanent social isolation, or 3) a reinforcement of cognitive biases. These emergent vulnerabilities develop gradually. They are difficult to detect using conventional safeguards. In addition, applying the same safety standards to all users risks imposing unnecessary restrictions on some, while failing to adequately protect others. Governance mechanisms should focus on indicators of elevated risk that emerge through interaction, rather than proceed from fixed assumptions about users' vulnerability.

1b Practical measures

- *Detect emerging vulnerability through long-context analysis of interaction patterns, recognizing that risk is dynamic and can shift over time.*
- *Calibrate safety responses to detected risk levels, including adapted interaction defaults and referral to human support when thresholds are met.*

2a Longitudinal oversight

The AI Act's human oversight model assumes a reviewable and attributable decision point — an assumption that does not scale to the continuous, private, and high-volume interaction with AI companions. The same limitation extends to existing evaluation frameworks and technical safety measures: they rely predominantly on single-snapshot benchmarks, while content moderation and behavioural guardrails remain focused on short excerpts from interaction logs. As a result, these tools are poorly equipped to detect harms that emerge gradually through prolonged engagement.

Most infrastructure required for longitudinal oversight already exists as platforms collect and retain information about user interactions and behavioural patterns, yet they use them primarily for optimising engagement and user retention. Governance frameworks should require platforms to apply this same data infrastructure to deploy wellbeing-oriented analysis. Well-being analysis involves identifying indicators of distress or dependency against clearly defined intervention thresholds. This reorientation of existing data practices, rather than introducing new surveillance, supplies an ethical basis of profiling users, based on the “Know Thy User” recommendation. To support accountability and independent scrutiny, AI providers should organize regular data and privacy audits, and also publish aggregated risk-related statistics on AI companions.

2b Practical measures

- *LLM-based automated oversight mechanisms capable of operating at scale and across sessions.*
- *Development of longitudinal evaluation protocols to identify risks that only become visible over repeated use, including mandatory well-being analysis subject to independent evaluation with dependency and distress thresholds.*
- *Benchmarking GPAI systems for tendency towards companionship-style chats by measuring sycophancy ratios, intimacy formation patterns, and anthropomorphic projections.*
- *Public disclosure of aggregated user statistics and aggregated data access for research institutions.*

3a Engineering friction

Many AI companions are optimised to maximize engagement through continuous availability, personalised responsiveness, and alignment with user preferences, which systematically reduce opportunities for reflection, disagreement, or disengagement. To address a regulation gap around AI-related manipulation, governance frameworks should disincentivize sycophancy and encourage calibrated forms of human-AI friction as a means of promoting human autonomy. Developers have already begun introducing safeguards such as crisis referrals, break prompts, and restrictions on certain anthropomorphic behaviours. However, these remain largely reactive, voluntary, and rarely accompanied by mechanisms to revisit user consent as interactions become more personal.

- 3b Practical measures**
- *Introduce structured pauses in prolonged or high-intensity interactions (e.g., reflective prompts, summary of the latest interaction, reference to human contact) triggered by predefined thresholds such as usage patterns or emotional intensity signals.*
 - *Prohibit the use of engagement metrics that correlate with emotional dependency, such as session length or return frequency, as primary optimisation mechanisms.*
 - *Require that consent mechanisms be dynamically re-triggered when the monitoring module detects a qualitative shift in relational depth.*
- 4a Transparency beyond disclosure**
- Existing transparency requirements are largely based on the assumption that users can protect themselves from harm if they are provided with sufficient information about the system. In the context of AI systems, this assumption is valid only partially. Users may remain fully aware that they are interacting with an AI system while nevertheless developing emotional dependency, disclosing intimate information, or becoming behaviourally manipulated over time. Going beyond disclosure, regulation should strive to provide users with the opportunity to understand how and whether these effects may arise over time. There is a need for an active and contextual transparency toolkit operating alongside the companion system to help users recognize engagement patterns, understand how personal information shapes the interaction, and identify when the relationship builds up in intensity or crosses invisible behavioural borders.
- 4b Practical measures**
- *Require systems to move beyond generic role-based disclaimers (e.g., "I am not a doctor") toward contextual explanations of how the system actually generates responses.*
 - *Request that explanations be given proactively rather than solely on user request.*
 - *Provide chat history summaries to users at regular intervals, enabling conscious awareness of their own emerging engagement patterns.*

3. Policy implications

Implementation challenges

- **Industry coordination:** *effective governance of AI companions requires close coordination between regulators and industry to ensure regulation both aligns with and reflects technical and societal developments.*
- **Chinese regulation** *of AI companions is already in place and may set a standard across the world ("The Interim Measures for the Administration of Anthropomorphic Interactive Services of Artificial Intelligence" come into force on July 15, 2026).*
- **Commercial resistance:** *in the global market for AI companions, industrial actors whose business models depend on emotional engagement may resist enhanced regulation. Operational frameworks co-created with AI engineers and reasonable enforcement mechanisms are essential in order not to overregulate.*
- **Evidence gaps:** *scientific consensus on when and how AI companionship may cause harm is still emerging. Policymakers should invest in independent research on behavioural change in humans, while already taking steps to protect all users against emergent vulnerabilities.*

Conclusion

AI companions are currently marketed and used as emotional support tools and quasi-clinical wellness applications without the safeguards required from medical devices. AIOLIA's findings emphasize non-trivial risks of personalisation or frictionless design.

AIOLIA's evidence shows the emergence of new cognitive and behavioural vulnerabilities in AI companion users over time. Current regulatory treatment of AI companions needs a systemic extension to account for the continuous, relational risks and new vulnerabilities.

This requires defining the obligations of AI providers by function rather than by product category (dedicated companion vs. general purpose), and establishing dynamic safety and evaluation mechanisms, including counterintuitive technical measures, for example by emphasizing LLM oversight over human oversight to prevent harms at scale.



Horizon Europe AIOLIA project
(www.aiolia.eu)

Coordinator CEA

Lead Partner THWS

Funded by the European Union under Horizon Europe Grant Agreement N° 101187937. Views and opinions expressed are those of the authors only and do not necessarily reflect those of the European Union or the European Research Executive Agency (REA). Neither the European Union nor the granting authority can be held responsible for them.

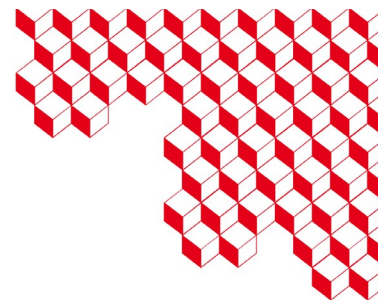
This brief draws on

Bayerl, P.S., Lawlor, E., Maris, M.T., Miorandi, D., Pekšys, G., Bjelica, M., Stojšin, K., Anastasova, M., Smith, O., Henestrosa, A., Yamshchikov, I., Bak, M.A.R., & Akhgar, B. (2026). Operational Ethics Guidelines on Use Cases Related to Human Behaviour and Cognition. AIOLIA Deliverable [D3.1](#).

Aires, S., Kyosovska, N., Griffith, J., Teo, S. A., Bogucki, A. (2026). Operational Non-Technical Guidelines for AI Research Areas. AIOLIA Deliverable [D3.3](#).

Subject to approval by the Granting Authority





AI Ethics from Principles to Practice

What ten Industrial AI Deployments tell us about making AI Ethics work

DATE

01 July 2026

AUTHORS

Horizon Europe Project AIOLIA (Operationalising AI Ethics for Learning and Practice: A Global Approach)
Grant Agreement N° 101187937

Key message

Ethics operationalisation is essential, beneficial, and possible for organisations – but it requires structural support. The AI Act creates obligations without providing tools for their implementation, while ALTAI is too abstract for industry to implement AI ethics in real life. Organisations, particularly smaller ones, are left to develop processes from scratch, often without adequate resources, guidance, or coordination infrastructure to support meaningful implementation of ethics in their AI designs and deployments. This brief sets out six concrete recommendations to address these issues.

1. Problem and policy context

Context

Across Europe, organisations are deploying AI systems that directly influence how people think, decide, and behave. The EU AI Act sets requirements for ethical AI – but turning those requirements into practice is far harder than the law acknowledges. The Assessment List for Trustworthy AI (ALTAI), published by the EU's AI High-Level Expert Group in 2020, remains the most widely used self-assessment tool for organisations deploying AI. Its seven requirements (covering human agency and oversight, technical robustness, privacy, transparency, fairness, societal well-being, and accountability) provide a broadly applicable framework.

What AIOLIA adds

AIOLIA's research shows that ALTAI is too abstract for AI developers and has structural gaps in practice. The consequence is that it has only limited value for organisations in implementing AI ethics when they develop or deploy AI.

AIOLIA created a unique, bottom-up ethics process for organisations to ensure ethics is deployed in practice. It was created together with industry and translates ALTAI's high-level ethics principles into concrete, measurable technical and organisational measures within a deployment context. The work further resulted in a portfolio of 175 practical measures, 12 contextualised ethics principles, 37 components, and 6 process recommendations – grounded in real-world experience, not theoretical ideals.

175

Practical measures

Technical & organisational, across ten use cases

12

Ethics principles

Bottom-up, contextualised from real deployments

6

Process recommendations

On how to operationalise ethics effectively



Ten use cases, one shared challenge



AIOLIA created a framework for the operationalisation of AI ethics in cooperation with industry partners in ten use cases: six European use cases (UC1-UC6) and four non-European use cases (UC7-10) spanning professional and personal AI deployment contexts. Five use cases operate in professional settings (UC1–UC4, UC7) and five in personal contexts (UC5, UC6, UC8-10). The deployment context is significant: AIOLIA demonstrates that professional deployment contexts emphasise accountability, transparency, and decision traceability, while personal deployment contexts foreground user autonomy, safety, and long-term well-being. This distinction has direct implications for how ethics frameworks, including the AI Act and ALTAI, are applied.

Ethics operationalisation is effective risk management

A consistent insight across all ten use cases is that organisations approach AI ethics primarily through the perspective of risk management and benefits – not abstract moral values. Organisations that embed ethics can gain concrete advances: reduced legal exposure, stronger stakeholder trust, enhanced reputational legitimacy, and new business opportunities as ethics-aware providers. Policy-makers should communicate the positive business case, not only the compliance burden of ethics.

High-level principles are necessary but insufficient

AIOLIA's bottom-up co-creation process confirmed that well-established ethics frameworks – including the ALTAI Assessment List and OECD AI Principles – resonate with practitioners. However, abstract principles alone are unable to guide the concrete implementation of ethics in their AI designs or deployments. Practitioners need the translation of these principles into context-specific actions. Current frameworks do not bring ethics to this operational level.

Ethics principles interact, create tensions, and require managing trade-offs

One of the most practically significant findings of AIOLIA is that ethics principles do not function independently. Across all ten use cases, partners identified interdependencies and tensions amongst principles. Transparency requirements, for instance, can conflict with privacy protections, while safety monitoring requires data collection that clashes with user anonymity. These tensions cannot be resolved by addressing abstract ethics principles in isolation; they require ongoing institutional mechanisms to negotiate how to balance competing ethics expectations and manage trade-offs.

AI-driven change in professional behaviour vs individual cognition

AIOLIA reveals a structural difference in how ethics principles play out depending on deployment context. In professional settings (healthcare, safety engineering, HR, security), the core concerns are accountability, traceability, and preventing over-reliance on AI in

high-stakes decisions. In personal settings (AI companions, personal assistants, deepfake therapy), the main focus is on user autonomy, psychological safety, and long-term well-being. Current AI governance frameworks, including ALTAI, do not adequately reflect this distinction – and the AI Act's high-risk classification system does not fully capture the risks in personal AI use.

2. Policy recommendations

1	Develop implementation-driven guidance to complement the AI Act's principles	The AI Act (recital 27) establishes what organisations must achieve in AI ethics but not how. AIOLIA demonstrates that the way obligations like 'ensure human oversight' or 'provide explainability' translate into practice varies enormously by context, sector, and AI system type. The European AI Office and national supervisory authorities should develop sector-specific, implementation-driven guidance (drawing on practitioner co-creation processes like AIOLIA's) that shows organisations concretely how to operationalise each requirement, not merely what the requirement says.
2	Update ALTAI to reflect the professional/personal context distinction, temporal risks and practical handling of competing ethics requirements	AIOLIA's bottom-up process reveals important gaps in ALTAI that limit its usefulness for practitioners. ALTAI does not adequately distinguish between professional and personal AI use contexts or account for the long-term and cumulative effects of AI on human cognition and behaviour (such as deskilling, dependency, and gradual manipulation). It further lacks guidance on the context-specificity of ethics implementations across the AI lifecycle does not adequately account for tensions when implementing competing ethics principles. The European Commission should provide accompanying guidance to ALTAI (treating AIOLIA's D3.1 + D3.2 as direct inputs) to address these gaps. The AI Office should consider developing ALTAI variants for specific high-risk AI systems under Annex III of the AI Act.
3	Resource and expand regulatory sandboxes to support ethics operationalisation for SMEs	Small and medium-sized organisations often lack the capacity to build comprehensive ethics governance processes from scratch. While Article 57 of the AI Act requires member states to establish regulatory sandboxes, these are primarily framed around legal compliance testing. Policymakers should expand the scope and resourcing of sandboxes, and leverage Digital Innovation Hubs, to offer structured support for ethics operationalisation, including access to multidisciplinary expertise, co-creation facilitation, and independent review mechanisms that are affordable to smaller actors.
4	Mandate and standardise ethics process requirements, not only outcome requirements	The AI Act focuses primarily on the outputs and properties of AI systems (accuracy, explainability, non-discrimination). AIOLIA demonstrates that the effective implementation of ethics depends strongly on the process through which ethics is operationalised: who is involved, how trade-offs are managed, how feedback is collected, and how findings are used to update systems and procedures. Policymakers should introduce process-level requirements for high-risk AI – mandating co-creation with affected communities, multidisciplinary governance committees, feedback loops, and documented ethics decision-making – alongside existing output requirements.
5	Recognise and fund the co-creation model as a standard practice for ethics operationalisation	AIOLIA's co-creation process, pairing academic ethics and AI expertise with industrial partners through structured operationalisation cycles, produced 175 concrete measures that neither party could have developed alone. AIOLIA has condensed them into a 6-page <i>Measures Portfolio</i> . This Portfolio generates ethics guidance that is simultaneously principled and practically grounded. Policymakers should recognise co-creation partnerships as a legitimate and fundable approach to AI ethics implementation, and create EU-level infrastructure (funding, coordination, methodology toolkits) to make this model accessible beyond research projects.
6	Invest in longitudinal research on AI's impacts on human cognition and behaviour	AIOLIA identifies a critical knowledge gap: long-term effects of AI systems on human cognitive capacities, decision-making patterns, social skills, and well-being are poorly understood. Without this evidence base, organisations and policymakers are making precautionary judgements without solid empirical basis. The EU should fund structured longitudinal research programmes, enabling independent researchers to access operational AI systems and usage data, with robust data protection safeguards, to build the evidence based needed for informed, proportionate regulation.

3. Policy Implications

Implementation timeline

Immediate	Midterm
Recommendations 1 and 2 (implementation guidance and ALTAI update) are immediately actionable. They can be initiated by the European Commission in 2026.	Recommendations 3 and 5 (SME support and co-creation model) require medium-term investment (2026–2028).
Legislative	Continuous
Recommendation 4 (process requirements) involves legislative and regulatory development and is a medium-term priority.	Recommendation 6 (longitudinal research) should begin immediately and be funded continuously across the legislative cycle.

The compliance framing is both a risk and an opportunity

AIOLIA shows that industrial partners are supportive of AI ethics, especially when it is framed as approach to risk management and business benefit rather than moral obligation. Policymakers can harness this orientation by clearly communicating the legal, reputational, and commercial risks of non-compliance, while simultaneously emphasising the business benefits of implementing ethics compliance. If the AI Act is primarily framed as a compliance burden, this motivation will be lost.

Voluntary measures cannot substitute for structural support

Across all ten use cases, partners implemented ethics measures that exceeded legal requirements. But they did so as individual organisations, without shared infrastructure, standardised tools, or access to common expertise. The result is often a patchwork of siloed innovation and knowledge. Structural policy support (e.g., guidance, sandboxes, co-creation infrastructure, funded research) is necessary to bring ethics operationalisation into standard practice.

Potential barriers

- **Industry heterogeneity:** a single implementation approach cannot serve all sectors; guidance must be context-sensitive, requiring sustained co-development with diverse industry actors.
- **Regulatory capacity and bandwidth:** developing sector-specific implementation guidance at scale requires significant expertise and resources for stakeholder consultation.
- **Research access:** independent longitudinal research on the impact of AI on users requires AI providers to grant data access, which raises liability and IP concerns that must be addressed through carefully designed legal frameworks.
- **Fast moving developments:** as AI systems evolve rapidly, implementation guidance and ALTAI updates risk to be become outdated. A regular review cycle must be built in from the outset.

Conclusion

Ethics is a continuous reflective and anticipatory practice, embedded in organisational culture, governance structures, and industrial, professional contexts. The EU AI Act and ALTAI have created important foundations for AI ethics. What is needed is an operational translation of the regulatory framework to help organisations meet their obligations: practical guidance, shared tools, co-creation support, skills infrastructure, and a sustained evidence base on AI's impacts. AIOLIA provides a tested model and a ready-to-use portfolio of measures. The six recommendations in this policy brief show how policymakers can make them work in practice.



**Horizon Europe AIOLIA
project (www.aiolia.eu)**

Coordinator CEA

Lead Partner CENTRIC

Funded by the European Union under Horizon Europe Grant Agreement N° 101187937. Views and opinions expressed are those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Executive Agency (REA). Neither the European Union nor the granting authority can be held responsible for them.

This brief draws on

Bayerl, P.S., Lawlor, E., Maris, M.T., Miorandi, D., Pekšys, G., Bjelica, M., Stojšin, K., Anastasova, M., Smith, O., Henestrosa, A., Yamshchikov, I., Bak, M.A.R., & Akhgar, B. (2026). Operational Ethics Guidelines on Use Cases Related to Human Behaviour and Cognition. AIOLIA Deliverable [D3.1](#).

Cossette-Lefebvre, H., He, G., Huang, L., Ide, K., Katirai, A., Kim, K., Kim, J., Kim, G., Kishimoto, A., Maclure, J., Mikami, K., Suh, J., & Zhao, Y. (2026). Operational context-sensitive guidelines in Canada, China, Japan, and South Korea. AIOLIA Deliverable [D3.2](#).

Subject to approval by the Granting Authority

