

JUST-AI JEAN MONNET CENTRE OF EXCELLENCE

CLOSING CONFERENCE

The Future(s) of HumA(I)nity: Human Agency, Agentic AI and Agentified AI Governance

Brussels, 29-30 October 2026

Who is the agent?

Within this conceptual framework, AI systems may be understood as mediating freedom in distinct and normatively significant ways. Rather than conceptualising such systems as either neutral tools or autonomous agents, we propose to examine how they participate in shaping the conditions under which freedom is exercised. In this regard, a useful distinction can be drawn between freedom-supporting and freedom-reducing forms of mediation. The former includes practices that enhance users' capacity for reflection, articulation, and informed decision-making, such as clarifying arguments, expanding interpretive options, and improving conceptual or textual expression within user-defined goals. The latter refers to forms of influence that constrain, distort, or bypass deliberation, thereby undermining the user's capacity for self-determination.

Crucially, both modalities involve some form of influence over action; the normative difference does not lie in the presence or absence of mediation, but in the conditions under which such mediation operates. These conditions may be articulated through three interrelated criteria: transparency, understood as the possibility for users to recognise and reflect upon the system's role in shaping outcomes; revisability, understood as the user's capacity to contest, modify, or override system outputs; and alignment, understood as the degree to which system behaviour remains responsive to the user's own ends rather than substituting them. Taken together, these criteria offer a way of operationalising the concept of freedom in AI-mediated environments without reducing it either to non-interference or to an unrestricted notion of technological co-production.

The JUST-AI Jean Monnet Centre of Excellence closing conference asks a fundamental question: *whose* agency should align with and deliver fairness? Is it a human actor - and if so, which one? Is it an artificial, agency-capable artefact or software system? Or is it an

institutional actor, such as a public or private agency tasked with monitoring and governing human-AI interactions?

Revisiting fairness in light (and in service) of humancentrism

Building on fairness as the central theme that has connected previous editions of the JUST-AI Jean Monnet Summer Colloquia, the closing conference of the JUST-AI Jean Monnet Centre of Excellence invites scholars at all stages of their careers to revisit the question of humancentrism.

Indeed, from the early stages of debates on AI ethics and regulation, humancentrism emerged as a guiding principle intended to provide an axiological framework for both ethical reflection and regulatory intervention. Although the concept has been interpreted in various ways, its underlying aspiration has generally been to steer technological development in a manner that preserves rather than undermines humanity, understood by ethicists and legal scholars alike through the lens of human dignity as a defining ontological and normative principle.

The proposition that AI technologies ought to serve and preserve humanity is at once compelling, intuitively appealing, and conceptually demanding. In his robot-centred imaginaries, Asimov articulated the imperative that humans should not be harmed. Recasting this intuition of “doing no harm” within contemporary debates has led scholars in digital humanities, fundamental rights theory, and dignitarian approaches to observe that, while “humanity” is too broad and complex a notion to function as a precise regulatory object, one aspect of humanity appears particularly exposed to the challenges posed by AI: human agency. The central concern, in this respect, is that AI systems should not deprive human beings of their capacity to act, decide, and shape their own lives. Yet agency is itself a multifaceted and contested concept. For the purposes of this conference, it will be approached through three interrelated dimensions: human agency, technological agency, and institutional agency.

Human agency: obsolete, transformed and... still free?

Arguably, the diverse and increasingly rich strands that make up contemporary AI scholarship - from substantive regulation, fundamental rights protection, and the rule of law to accountability and liability - are, in one way or another, concerned with the theoretical question of uncoerced human agency. The normative, political, and legal imperative that AI technologies should support rather than restrict human freedom underpins much of the social-scientific literature on AI, albeit through different conceptions of what such freedom entails.

These conceptions of uncoerced agency include, among others, the capacity for self-determination; the effective and unrestricted exercise of fundamental rights; freedom from being nudged or manipulated into making particular commercial or non-commercial choices; the entitlement to understand how AI systems operate and affect individuals; and the ability to allocate accountability and liability in an appropriate and meaningful manner.

These predominantly legal approaches tend to share a common feature: they conceive of human agency primarily in negative terms, as a form of “freedom from” interference. The

underlying assumption is that human potential flourishes most fully when protected - through legal safeguards such as rights, duties, and institutional guarantees - from intrusive, manipulative, or otherwise constraining AI technologies.

Agency-apt AI

As a feature of AI technologies, agency is commonly understood in two broad ways: either as an extension of human agency or as something increasingly distinct from it. As technological development has accelerated, we have witnessed the emergence of AI systems that appear, as it were, progressively emancipated from direct human tutelage. These systems exhibit degrees of autonomy that no longer require humans to predefine every objective they pursue, nor do they always permit effective human control or oversight over the means by which those objectives are achieved.

Agentic AI, often presented as a significant step toward Artificial General Intelligence (AGI), exemplifies this trend. Such systems are increasingly capable of identifying, refining, and pursuing goals on behalf of human actors, and in some cases with a degree of independence from them. In many everyday contexts, these capabilities may be viewed as usefully complementing and augmenting human agency. Yet they also raise important normative questions, as the levels of autonomy displayed by such systems challenge assumptions embedded in many legal and ethical frameworks, particularly those grounded in the principle of human dignity and its commitment to human self-determination.

While the delegation of decision-making to AI systems may be relatively unproblematic in certain domains, there are contexts in which the minimization or exclusion of meaningful human agency can have profound real-world consequences. Military applications of AI provide a particularly telling example. The increasing use of autonomous and semi-autonomous systems in contemporary warfare raises pressing questions regarding human control, accountability, responsibility, and the conditions under which decisions affecting human life may legitimately be delegated to technological agents.

Agentification of AI governance: institutional specialization (fissuring?) in the search for coherence

In discussions on AI, governance has often been closely associated with the EU AI Act (Regulation (EU) 2024/1689). The narrative is by now familiar: the AI Act establishes a risk-based regulatory framework designed to ensure compliance through *ex ante* and *ex post* conformity assessments, market surveillance mechanisms, and institutional oversight structures. Its overarching objective is to safeguard individuals and society while enabling the responsible development and deployment of AI technologies.

For the purposes of our closing conference, however, we adopt a broader understanding of AI governance. Rather than limiting governance to the institutional architecture established by the AI Act, we conceive of it as encompassing the wider range of regulatory, political, technical, and societal arrangements through which AI technologies are shaped, monitored, and governed. This broader perspective extends beyond questions of consumer safety and includes domains increasingly transformed by AI, such as sustainable development, intellectual property and artistic creation, urban governance and smart cities, scientific research, public administration, and critical infrastructures.

In a context in which the protection of human dignity remains a foundational normative commitment, yet AI systems appear increasingly difficult to scrutinize, control, and predict, questions of governance acquire renewed importance. We therefore invite contributions examining the design, operation, and legitimacy of governance models that regulate AI either as a distinct object of regulation or as a significant component of regulatory domains not traditionally conceived as dependent upon technological innovation, including environmental protection, sustainability, culture, urban infrastructures, and public services.

Particular attention will be given to the institutional dimensions of AI governance. A central question is whether contemporary developments reveal a trend toward the fragmentation - or "agentification" - of AI governance. Such a trend would be characterized by the proliferation of specialized regulatory bodies, standard-setting organizations, expert committees, technical authorities, and certification entities operating at national, regional, and international levels. Examples include standardization organizations such as CEN and CENELEC, technical bodies such as NIST, as well as the growing number of expert groups, advisory boards, and governance initiatives dedicated to particular aspects of AI development and deployment.

This raises a broader theoretical question concerning institutional agency: does the multiplication of governance actors enhance the effectiveness and legitimacy of AI regulation, or does it risk fragmenting regulatory authority and undermining the coherence of legal frameworks specifically designed to govern AI, such as the AI Act? More generally, how should responsibilities be distributed among the increasingly diverse constellation of public, private, and hybrid actors involved in governing AI technologies?

The JUST-AI JMCE Closing Conference: a platform for pluri- and interdisciplinary discussion

The JUST-AI JMCE Closing Conference aims to bring together researchers in various stages of career advancement (doc, post-doc, Assist. and Assoc. Prof) active in **IT**, the **social sciences** (sociology, political science, law and economy) and the **humanities** (philosophy, history, linguistics). While contributors are encouraged to submit proposals on **any topic** in line with the theme of the conference, the following topics may provide guidance on the types of presentations and discussions we aspire to have:

1. Agency, self-determination, fundamental rights

- Human dignity as cornerstone of AI regulations around the world
- Digital identities
- Transduction, co-participation or cocreation: examining the nature of human-AI interactions
- Agency of vulnerable individuals and groups
- Dark patterns, fundamental rights and consumer protection
- Manipulative technologies and the effectiveness of consent
- GDPR, AIA and the perils of LLms
- The feasibility and effectiveness of Fundamental Rights Risk Assessments
- Human control and oversight as fairness-enhancing principles
- Enhancing agency through procedural abilities: explainability as a procedural device
- Explainability and counterfactual reasoning
- Standards and best practices for data debiasing strategies
- Auditing and measuring biases
- The concept of collective AI-related harm

2. Agency-apt AI, agentic AI

- The power of AI: Palantir, Fable 5 and the threats to individual and State power
- Agentic AI and the role of humans in autonomously defining objectives
- The gaps in the regulatory capture of agentic AI
- The future of human control and oversight: between performativity and effectiveness
- AI in warfare: the perils of human control becoming obsolete
- AI in mass surveillance
- Automated or reasoned adherence to AI output: the art of making 'adequate' AI-assisted decisions
- The future of transparency: noble ideal or enforceable (legal) requirement?
- The reframing of trustworthy AI: should we trust human-independent AI technologies?
- The concept of accuracy and the challenge of reliably performing accuracy checks

3. Agentification of AI governance

- Expertise, political and normative powers of standardization bodies
- The role of experts in the implementation of AI regulatory frameworks
- The role of private governance in operationalizing transparency and fairness
- Market surveillance authorities and the effectiveness of post-market monitoring
- Shortcomings in ex-ante conformity assessment procedures under the AIA
- Public and private actors in the AIA's implementation: in the search for a coherent approach
- Rethinking the risk-taxonomy in the AIA: what does the AIA leave out?
- Administrative redress, out-of-court dispute resolution and judicial redress: trends across the EU's digital legislation
- Public security as justification to mass surveillance
- The principle of proportionality in the exercise of administrative discretion
- The duties to safeguards procedural rights in AI-assisted investigative and dispute-resolution procedures
- Constitutional (re)framing of AI-assisted governance

Format of submissions & deadline

All submissions must be in **English** and include the following:

- **Cover page:** full name(s), affiliation, contact information of the author and title of the submission
- **Abstract** (separate page): 500 words (list of references is not mandatory)
- **Short CV** (separate page): one paragraph

These should be included in a **single Word or PDF file** and sent to: just.ai@uliege.be by **1 September 2026**.

Please use the following template for the subject of your email: **Submission_JUST-AI_2026_ClosingConf._Family Name_Title of Contribution**.

Review, acceptance and financial support

All submitted papers will undergo a blind peer-review process, performed by **at least two members of the JUST-AI team**: one from the [JUST-AI Governance Board](#), the other from the [JUST-AI Editorial Board](#). Pluri- and interdisciplinary submissions may be reviewed by one to two additional experts. The reviewers will assess the formulation, originality and relevance of

the selected topics as well as the quality and coherence of the main arguments (and the normative claims, if any) presented by the authors.

Ten to twelve abstracts will be selected to be presented in Brussels, before an interdisciplinary panel of experts (5-6).

The JUST-AI Jean Monnet Centre of Excellence will support the selected authors financially for their travel (round-trip plane/train tickets) and accommodation (1 night in Liège) costs.

What happens next?

After the closing conference, the speakers will be required to submit studies (10'000 words incl. footnotes) based on the presented abstracts. These studies will undergo additional reviews in view of publication in a **special issue of a specialized journal or review or an edited volume**.

The studies presented at the 2024 edition of the JUST-AI Summer Colloquium are published in a special issue of the [Cambridge Forum on AI: Law and Governance](#).

The studies presented at the 2025 edition of the JUST-AI Summer Colloquium are featured in an edited volume published by **Routledge**, in the [AI, Law and Society Series](#).

Information on the editorial requirements, submission deadlines and review process will be communicated to the authors subsequently.

Important dates

Deadline for submitting abstracts: **1 September 2026**

Notification of acceptance: **21 September 2026**

Dates of the Conference: **29-30 October 2026**

- *travel to Brussels is encouraged, though online presentations may be considered if the speakers are unable to attend in person.*

About the JUST-AI Jean Monnet Center of Excellence

Launched in 2023 under the European Commission's Jean Monnet Centre of Excellence label, the **JUST-AI JMCE** strives to become a Belgian and European 'hub of excellence' on issues pertaining to procedural fairness in the field of AI. With the participation of a network of Belgian and international experts on AI, and in view of reaching out to expert and non-expert audiences around the world, the activities launched under the Centre of Excellence label aim at promoting innovative, pluri- and interdisciplinary research, creating original and interactive educational content and raising awareness on the most salient current and future issues on procedural fairness in the era of AI. The Centre's activities will, in particular, be developed along two research axes: 1. AI liability and fair trial requirements and 2. regulation and procedural capabilities of litigants in AI-related disputes. To learn more about the actions of the JUST-AI JMCE, you are welcome to browse through our [website](#).

If you have questions pertaining to the Center of Excellence and/or the JUST-AI Closing Conference, please email us at just.ai@uliege.be.