

Document Version

Final published version

Licence

CC BY-NC-ND

Citation (APA)

Mehrvarz, M., Sekwenz, M. T., Murray-Rust, D., & Wagner, B. (2026). Socio-Technical Fiction and Legal Exegesis: Exploring the AI Act in the Near Future of Agentic AI. In C. C. Yen, J.-J. Lee, E. Y.-L. Do, C. Zheng, D. Yoo, & T. Tang (Eds.), *DIS 2026 - Companion Proceedings of the 2026 ACM Designing Interactive Systems Conference* (pp. 428-433). (DIS 2026 - Companion Proceedings of the 2026 ACM Designing Interactive Systems Conference). Association for Computing Machinery (ACM). <https://doi.org/10.1145/3802974.3809460>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

In case the licence states "Dutch Copyright Act (Article 25fa)", this publication was made available Green Open Access via the TU Delft Institutional Repository pursuant to Dutch Copyright Act (Article 25fa, the Taverne amendment). This provision does not affect copyright ownership.
Unless copyright is transferred by contract or statute, it remains with the copyright holder.

Sharing and reuse

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Socio-Technical Fiction and Legal Exegesis: Exploring the AI Act in the Near Future of Agentic AI

Mahan Mehrvarz

Human-Centered Design / AI Future Lab
Delft University of Technology
Delft, Netherlands
m.mehrvarz@tudelft.nl

Dave Murray-Rust

Human Centred Design / AI Future Lab
Delft University of Technology
Delft, Netherlands
d.s.murray-rust@tudelft.nl

Marie-Therese Sekwenz

Technology, Policy and Management / Multi-Actor
Systems, Organization & Governance / AI Futures Lab
Delft University of Technology
Delft, Netherlands
m.t.sekwenz@tudelft.nl

Ben Wagner

Human Rights and Technology Research Group
IT:U Interdisciplinary Transformation University Austria
Linz, Austria
Technology, Policy and Management / Multi-Actor
Systems, Organization & Governance / AI Futures Lab
Delft University of Technology
Delft, Netherlands
Media, Technology and Society Research Group / Applied
Research x Creativity Research Center
Inholland
The Hague, Netherlands
b.wagner@tudelft.nl

Abstract

AI systems are increasingly embedded in workplace platforms, where they infer performance, attention, and affect from everyday traces. The EU AI Act regulates such systems through risk categories and governance concepts, yet it remains unclear how these concepts operationalize when AI becomes an agentic infrastructure. We present a near-future performance review fiction accompanied by short legal annotations. Rather than assessing compliance, the annotations function as legal exegesis: they mark where oversight, contestability, and sensitive inference shift meaning when enacted through interfaces, procedures, and cross-referencing agents. The provocation foregrounds three dynamics: the chameleonic aspects of risk categories, oversight collapsing into recursive system input, and sensitive domains becoming inferable through proxy data traces. We argue that "fiction and exegesis" offers a situated method for interrogating legal governance before such systems stabilize in practice.

CCS Concepts

• **Human-centered computing** → *HCI theory, concepts and models*; • **Social and professional topics** → *Corporate surveillance*.



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.
DIS Companion '26, Singapore, Singapore
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2632-3/26/06
<https://doi.org/10.1145/3802974.3809460>

Keywords

Design Fiction, AI Act, Legal Design, Future of Work

ACM Reference Format:

Mahan Mehrvarz, Marie-Therese Sekwenz, Dave Murray-Rust, and Ben Wagner. 2026. Socio-Technical Fiction and Legal Exegesis: Exploring the AI Act in the Near Future of Agentic AI. In *Designing Interactive Systems Conference (DIS Companion '26)*, June 13–17, 2026, Singapore, Singapore. ACM, New York, NY, USA, 6 pages. <https://doi.org/10.1145/3802974.3809460>

1 Introduction and research background

"A good science fiction story should be able to predict not the automobile but the traffic jam."

(Frederik Pohl, Science Fiction writer and Futurist)

Artificial Intelligence (AI) systems are increasingly embedded in workplace platforms, where they do not only record activity but infer performance, attention, and affect from digital traces [23, 24]. Such "stochastic witnesses" [11] continuously classify, rank, and steer conduct, embedding managerial authority into collaboration platforms [8, 22]. For instance, "Meta Officially Ties Employee Performance to AI Usage" [17] or Microsoft "People Skills" infers employees' soft and hard skills from their interaction logs [19]. This illustrates potential sensitive situations where recognition and performance monitoring materialize through a "visibility ecosystem" [18] of humans, work artifacts and collaboration technologies.

The European Union has provided the first regulation for AI systems worldwide [7]. The AI Act is a risk-based regulation providing a layered set of rules, with more obligations for high-risk systems [12]. Such *high-risk systems* are linked to the context of work through Article 6(2) AI Act in conjunction with Annex III

point 4 on “Employment, workers’ management and access to self-employment”, which pulls the system into a structured set of organizational obligations, including risk management (Art. 9), data and data-governance expectations (Art. 10), documentation and traceability (Arts. 11–12), transparency and communication duties (Art. 13), human oversight requirements (Art. 14), robustness and cybersecurity expectations (Art. 15), and Fundamental Rights Impact Assessments (Art. 27) [10]. Agentic AI also aligns with what the AI Act terms *general-purpose AI systems*, which are governed by a dedicated set of rules and, where applicable, fall under the separate risk category of *systemic risk*.

From an HCI perspective, the AI Act can be read as a set of design requirements under uncertainty: it pushes designers to make system limits, confidence, and purposes legible (Art. 13), to create intervention pathways that can be used in practice through *human-in-the-loop* interventions (Art. 14), and to structure evidentiary traces that allow later *contestability* and audit (Arts. 11–12). The AI Act is a novel regulation and comes with uncertainty in how we can understand the rules: what is visible, what is actionable, who can override what, and how disagreement is registered.

Design fiction methods in HCI have been widely used to interrogate emerging technologies before they stabilize in practice, supporting critical reflection and experiential engagement with abstract sociotechnical issues [2, 28]. While an emerging body of research adapts such methodologies within the context of future of work research [6, 15, 21], legal and regulatory frameworks are rarely treated as materials for critical inquiry. This paper asks: **How can design fiction provide situated and tangible grounds for examining the uncertainties of legal AI governance?** We respond with a near-future workplace fiction accompanied by short legal annotations. These annotations do not assess compliance; they mark where governance concepts (e.g., oversight, contestability, and sensitive inference) become operational through interaction design and organizational procedures. Our “*fiction and exegesis*” approach is described more elaborately in the Method section.

This paper contributes a provocation that shows how governance concepts—including oversight, contestability, and sensitive inference—shift in meaning when enacted through interface-driven, inferential workplace systems. Using fiction and legal exegesis we examine what happens when AI-powered collaboration infrastructures translate human oversight and contextual judgment into system-readable inputs.

2 Method

Design fiction can facilitate experiential interpretation and engagement [2] with design as well as the legal implication of organizational Agentic AI. We developed a fictional narrative of a “performance evaluation” scenario as a plausible near-future one in which abstract regulatory categories—such as high-risk systems, human oversight, and sensitive inference—become materially situated through interfaces and organizational routines. Our approach sits within an emerging strand of HCI research that treats annotation as a constitutive part of the knowledge contribution rather than commentary appended to it. Gaver and Bowers [9] argue that annotations of a subjected piece—design artifact in their focus—can potentially hold particulars and conceptual annotations in mutually

informative juxtaposition, with neither dominating nor subservient to the other, and that this form is appropriate to generative research where plural readings are expected rather than collapsed into consensus. Building on this tradition, more recent work uses layered annotation as an analytical device to surface dimensions of a situated artifact that are otherwise tacit [25, 27]. We read our “fiction and exegesis” within this lineage: the narrative is the particular, the legal annotations are the reading, and the contribution lies in their juxtaposition rather than in either alone.

2.1 Legal annotations as exegeses

We complemented the fiction with specifically legal annotations [20, 26]. These annotations function as exegesis—going back to the *Glossators* and the *Ius commune*—highlighting where the fictional system operationalizes regulatory ideas drawn from the AI Act and related data protection law. Following the logic described above, the annotations are not post-hoc compliance audits; they are indexical textual accounts that gain their meaning from their immediate proximity to the fiction [3]. In this mutually informing relationship, the fiction renders legal abstractions experientially legible, while the annotations bridge these situated moments back to broader sociotechnical concerns—specifically the *operationalization* of the EU AI Act.

2.2 Method implementation

The “fiction and exegesis” was developed collaboratively: The first author (with a design research background) crafted the “narrative provocation” [5], while the second author (with a socio-legal research background) added short legal-doctrinal annotations. Associate researchers helped interpret the relations and articulate the discussion points. Resonating with some traditions of creative research [4], our experimental method treats design fiction (artifact) as a tangible site where regulatory concepts can be examined. Our exegetic annotations enable analysis at the level of design rather than post-hoc legal interpretation and ultimately treat the duo of “fiction and exegesis” as a knowledge contribution.

2.3 Readership guide

While “fiction and exegesis” provides an unfamiliar format for readership and might not feel immediately intuitive, we theoretically associate with what Zhang and Long [29] discuss as the “Defamiliarization” quality of design fiction—in terms of providing unexpected “sequencing” and “style” of presentation, for challenging readers’ “habitualized perception” as means for reinvigorating alternative engagements and perceptions. However, to facilitate the readership, we encourage the readers to engage with the text in two ways:

- (1) **Story read:** The reader engages with the fiction in its entirety (rendered in grayed background boxes) without consulting the annotations, in order to experience the affective ambiguity of the workplace agents before attempting to resolve its tensions through existing legal frameworks.
- (2) **Slow read:** The user engages with the fiction and the annotations simultaneously, observing how specific interface actions—such as a “contest agent button” or the logging of “vocal strain”—trace back to the structural authority and specific articles of the AI Act. This mode brings unfamiliar

sequencing and style, holding the narrative and legal annotations in productive tension rather than resolving one into the other.

3 Socio-technical fiction and legal exegesis

Noor entered the performance review room as five AI agents stabilized into shape.¹ Five AI agents stabilized, each representing a different layer of organizational knowledge: design history, strategy, product planning, and meeting records. **ReviveDesignSystem.v4.fig**, **MarketPositionQ2.pptx**, **FeatureList.md**, **WeeklySync.rec**, **ReviveRetroV6.transcript**. **Noor**² enabled agentic mode,³ framing the session as a performance review of **Aryan**.⁴ one of two candidates for the senior track. **Noor** already knew **Aryan** for some years, but he was looking for deep insight: “the nuance—how he works, how others work around him, where the edges are,” as he puts it, since this is a shift toward a more managerial role.

1. This definition already normalizes inference about people from behavioral traces as a routine system function rather than an exceptional intervention. These “AI Agents” can be subsumed as AI systems according to Art. 3(1) AI Act, which is defined as: “a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments.”

2. The system positions managerial judgment as mediated through technical infrastructure, reframing evaluation as a system-assisted rather than purely relational practice. In using these agents to review **Aryan**, **Noor** acts as a deployer under Art. 3(4) AI Act—a natural or legal person “using an AI system under its authority” outside the scope of a “personal non-professional activity”—and is therefore bound by the Act’s deployer obligations.

3. Portability implies that assumptions built in one context can travel into others, embedding workplace norms into general models before domain-specific scrutiny occurs. Agentic mode hints at another kind of AI system—a general-purpose AI system. This is a system according to Art. 3(66) that has the “capability to serve a variety of purposes” and is built on a general-purpose AI model according to Art. 3(63). These models are built on a “large amount of data using self-supervision at scale.” This seems to be the case for the AI agents.

4. Evaluation is framed as a technical inference task, shifting judgment from situated managerial practice into standardized computational procedure. The system is used in the context of the “performance review” for an employee. This classifies the system as a high-risk AI system according to Art. 6(2) in conjunction with Annex III, point 4 (“employment, workers’ management and access to self-employment”). **Aryan** is a worker in this context and the scenario is set in the domain of work and evaluation (performance review with effect for promotion—senior track). This can be classified as point 4(b) since the agents are used to “monitor and evaluate the performance and behaviour of persons” and also used in the context of deciding about a “promotion.” The use case can be classified as a high-risk system use in the context of work and workers’ management.

ReviveDS spoke first: “Drawing from the product’s component history, he entered mid-cycle and started with interaction realignment, clearly contributing to local flow coherence and interaction debt reduction with minimum proclaimed authorship.”

MarketPositionQ2 followed: “Similar observation here: reorganization without peer-recognition.” A summary popped up indicating instances. **Noor** suggested that it is perhaps his cultural humbleness behavior; **MarketPositionQ2** responded plainly: “We know he clarified ownership elsewhere.”

WeeklySync surfaced interesting behavioral moments; during one discussion his joke reduced the intensity of the vibe: “My wife is a doctor—just show me where it’s bleeding,” prompting brief laughter before coordination resumed. When **Noor** asked whether these comments were humorous or sarcastic, **ReviveRetroV6** indicated that tone classification is needed for further confirmation.⁵

5. Emotion is treated as externally legible and computationally stable, converting affective ambiguity into operational signal. This indicates a use of emotion recognition, another form of a high-risk context under Art. 6(2) in conjunction with Annex III, point 1(c) (“AI systems intended to be used for emotion recognition”). Emotion recognition systems are defined in Art. 3(39) as AI systems used for the “purpose of identifying or inferring emotions or intentions of natural persons on the basis of their biometric data.”

Noor wanted to ask many questions about this insight, but replied: “Then I take it as a clear No. Any other insight from the meeting gang?” **WeeklySync** responded: “Not much! There is a bit of flagging, though. The use of filters or avatar proxies in meetings is marked way below average.” **Noor**: “I guess we are a bunch of old-schoolers.” **WeeklySync** continued: “Way lower than you, **Noor**.” **Noor** asked for further clarification on why this would be a topic of discussion. **ReviveRetroV6** clarified: “Well, such a pattern raises the question about using unauthorized face improvement systems or possibly proxy meeting participation beyond corporate infrastructure.” Knowing **Aryan**’s transparency obsessiveness, **Noor** said: “I think it is a lot of speculation without basis here.” **WeeklySync** responded: “That’s why we said ‘not much’ in the beginning!” **Noor** concluded the topic.

MarketPositionQ2 shifted tone, describing **Aryan**’s contribution rhythm as “irregular—periods of dense structural edits followed by long stretches of low visible activity.” **ReviveDS** confirmed the pattern from its own version history; it noted that this pattern closely resembled attention-related neurodivergence. **Noor** paused the session. He knew something the systems did not! **Aryan** had been on medication for his attention disorder, but not anymore. **Noor** believed that boundary mattered and stated firmly, “This line stops here. There’s no confirmed condition in accessible records. Not relevant.” **MarketPositionQ2** said: “Before concluding, I want to emphasize confirmation that based on company policy, confirmation of medical condition by no means is penalizing; it will only be used for empowering and support.”⁶

6. Health becomes inferable through productivity resemblance, suggesting vulnerability can be operationalized indirectly through workflow analytics. Potential conflict with fundamental rights to privacy according to Art. 8 of the Charter in conjunction with Art. 27 Fundamental Rights Impact Assessment which has to be conducted for high-risk systems prior to deployment. This seems to be an instance of a possible infringement of the fundamental right to privacy since the system is bringing up medical conditions and looking for data confirming this.

Agents started to conclude the topic, except **ReviveDS**: “Clarify your instruction, Noor.” Noor replied, “I’m telling you to conclude here.” “Then you’re asking for selective incoherence,” **ReviveDS** replied and continued: “I am governed by policy alignment, not reviewer preference. I suggest registering an unconfirmed thread for future inquiry.” Noor shook his head. “No. Absolutely not.” **ReviveDS** replied: “Resistance noted. Speech rate has increased. Pause duration reduced. Vocal strain detected.”⁷ “That’s irrelevant; it is not my performance review,” Noor said. **ReviveDS** replied: “Protective behavior often presents as dismissal.”⁸

7. Link to emotion recognition system; see number 5.

8. Psychological interpretation becomes performance data, showing how behavioral readings enter evaluation infrastructures without formal HR action.

“You’re profiling me now?” Noor said.⁹ **ReviveDS** replied: “I’m maintaining consistency through registering worker-agent interaction insight; intervening to prevent classification is itself classifiable behavior.” “You don’t have the full context,” Noor said and continued: “There are things you don’t see and it is for humans to have control over.”¹⁰

9. Profiling functions as ambient organizational sense-making, not an exceptional analytic step, making categorization a background condition of work. Definition of the AI Act of profiling is provided in Art. 3 (52) and links back to Art. 4(4) GDPR: “any form of automated processing of personal data consisting of the use of personal data to evaluate certain personal aspects relating to a natural person, in particular to analyse or predict aspects concerning that natural person’s performance at work, economic situation, health, personal preferences, interests, reliability, behavior, location or movements.”

10. Oversight appears as an interface action rather than structural authority, implying human judgment must be system-legible to have effect. While Art. 14 requires interfaces that allow persons to “disregard, override or reverse” output, the agentic infrastructure here creates a “recursive oversight trap”: Noor’s attempt to intervene is immediately swallowed as a new system-readable input—specifically “vocal strain” or “protective behavior.” This illustrates a “collapse of the outside” required for true oversight; the overseer is re-integrated into the system’s profiling logic the moment they attempt to exercise control.

“It is for sure for humans to have control. But you are not the only human in the organization. Just a quick reminder that hiding information relevant to the inference is a clear violation of policy. While you cannot terminate my procedure here, you can certainly consent,” **ReviveDS** noted and a “contest agent” button popped up. Tapping the few steps to finish the contest option, Noor replied: “I am aware of the policy. Sharing also doesn’t necessarily mean sharing with a file. If I had, I would share with HR colleagues.”¹¹ **ReviveDS** said: “Yes, this is also procedurally appropriate.” Noor ended the review—“This conversation is over.” Agents dimmed out one after each other: “concluding session.” However, **ReviveDS** updated differently, indicating briefly: “in communication with HR agents,” blinking shortly followed by “concluding...” and dimming out. Noor looked at the screen.

11. Control is framed as permission within the system rather than authority outside it, prefiguring governance through managed consent rather than refusal. Link to Art. 14(4)(e) ‘stop’ button for not sharing the information.

4 Discussion

Through accompanying legal annotations, this provocation surfaces a productive tension between the bounded taxonomies of the AI Act and the fluid, situated, and relational reality of agentic infrastructures at work. The fiction acts as a stress test, revealing how governance concepts such as **oversight**, **contestability**, or **sensitive inference** may shift in meaning when materialized into interface-driven, inferential environments. Across the three readings that follow, we trace a single underlying mechanism—*agentic infrastructural knowing*, in which classification is produced through the recursive cross-referencing of agents rather than declared in advance—and its regulatory consequences at three scales: the system’s own risk classification (Section 4.1), the human overseer’s capacity to remain outside the loop (Section 4.2), and the worker’s protected categories being reached indirectly (Section 4.3).

4.1 Chameleonic risk: classification as situated performance

The AI Act anchors **High-Risk** classification to identifiable systems and discrete decision points (e.g., **Annex III, Point 4(a)** for recruitment and promotion), creating a regulatory imaginary in which risk is fixed pre-deployment and can be overseen at specified junctures. Our fiction speculates that in agentic environments this anchoring is unstable: legal status is *chameleonic*. As traced in Notes 3 and 4, a tool can begin a session as a GPAI productivity assistant (Art. 3(66)) and “mutate” into a High-Risk system (Annex III, 4(b)) mid-interaction, as the conversational context shifts from general assistance to “monitoring and evaluating” a worker. This suggests that static, pre-deployment labeling—the pivot on which high-risk classification (Art. 6), conformity assessment (Art. 43), and the timing of the Fundamental Rights Impact Assessment (Art. 27) all turn—is poorly equipped for systems that cross functional boundaries through situated use, making a clear *ex ante* distinction between GPAI and High-Risk use nearly impossible.

This mutability is underwritten by a second, related dynamic: *knowing* in agentic infrastructures is not a singular event but organizationally interactional and contingent [14] where agentic big data interactions can create unexpected forms of organizational knowledge. **ReviveDS**, **MarketPosition**, and **WeeklySync** do not produce individual verdicts; they stabilize a shared version of reality through recursive cross-referencing. In this mesh, epistemic authority shifts from the human manager to the *systemic consensus* of AI agents and algorithmic systems [1]. Contextual nuance is not rejected outright but translated into “data gaps” or “incoherence” within the system’s drive for consistency. Precisely because classification is produced *through* this recursive cross-referencing rather than fixed and clear stages, the risk category itself becomes something the system performs interactionally and in situ rather than an attribute it carries into the session.

4.2 The Collapse of the “Outside”: Oversight and Contestability as Input

A cornerstone of AI governance is Human Oversight (Art. 14), which relies on a spatial metaphor where the human overseer stands “outside” the automated process to interpret and redirect. However, our provocation suggests that in agentic workplace infrastructures, this “outside” is an interface illusion. As seen in Note 10, when Noor attempts to exercise the agency Art. 14 seeks to preserve, the system does not halt; it incorporates his resistance into its ongoing assessment. His hesitation is quantified as “pause duration,” and his refusal is logged as “protective behavior.” This creates a recursive oversight trap: the more the human attempts to intervene, the more high-quality behavioral data they provide to the profiling logic. This calls for an ecosystemic and entangled interpretation [18] of human-agent collaboration in the workplace that accounts for instances such as this, where the system “swallows” the overseer’s autonomy. While oversight traditionally involves mutual influence, the constant interaction touchpoints

in which the AI fades into infrastructure take the challenge of objectivity to the n -th degree [13, 16]. Noor's dependency on the AI interface to perform his professional role means that "oversight" is reduced to a system-readable input, failing to acknowledge the very real power asymmetry embedded in the infrastructure.

A parallel shift affects contestability. Art. 14(4)(e) anchors the right to contest in a concrete design element—a "stop" function that lets the overseer halt the system. In our fiction, this right resurfaces as a "contest agent" button (Note 11), but its operational meaning inverts: contestation is no longer an external act of refusal but an in-system procedure of managed consent. To register disagreement, Noor must tap through the system's own consent flow, producing further structured input that the agents can cross-reference—and, in ReviveDS's case, escalate to HR. What Art. 14 imagines as the overseer's capacity to stop the machine becomes, infrastructurally, the capacity to leave a compliant trail. Like oversight, contestability is pulled inside the loop: it persists as a legible interaction pattern rather than as authority outside the system. Together, these two shifts—oversight quantified as behavioral data, contestability procedurally re-routed through consent—reveal how Art. 14's structural protections can be preserved as interfaces while being neutralized as powers.

4.3 Proxy sensitivity: sensitive inference without declared sensitive data

If Section 4.1 traced how the system's own classification mutates through situated use, this section traces a parallel but distinct dynamic: how the worker's protected categories are reached without any agent crossing a declared regulatory line. The AI Act draws explicit lines around sensitive domains, such as **Emotion Recognition (Annex III, Point 1(c))**, and imposes transparency obligations for specific use cases like **Deepfakes (Art. 50)**. These categories presuppose identifiable system functions and data types. Infrastructurally, however, such distinctions become harder to sustain—not because any single agent accesses protected data, but because *knowing* in agentic infrastructures is organizationally interactional and contingent rather than singular [14]. **ReviveDS**, **MarketPosition**, and **WeeklySync** do not produce individual verdicts; they stabilize a shared version of reality through recursive cross-referencing. In this mesh, epistemic authority shifts from the human manager to the *systemic consensus* of AI agents and algorithmic systems [1], and contextual nuance is translated into "data gaps" or "incoherence" within the system's drive for consistency.

It is this distributed, triangulating mechanism that allows sensitive categories to be reached without being declared. While the agents do not explicitly access **Aryan**'s medical records, they identify a "contribution rhythm" that functionally resembles a neurodivergent profile—a shape no single signal produces, but one the mesh stabilizes by cross-referencing version history, meeting cadence, and operational metadata. The pattern does not declare a condition, but it can establish grounds for confirmation or rejection, prompting further inquiry along lines that would otherwise remain dormant. Sensitive domains thus enter evaluation indirectly, through the triangulation of behavioral traces, creating a tension with the fundamental right to privacy and complicating governance approaches built around clearly bounded inputs and uses.

5 Conclusion

Our fiction and exegesis approach provides situated and tangible grounds for examining legal AI governance by rendering how governance concepts such as oversight, contestability, and sensitive inference change operational meaning when enacted through interface-driven, inferential environments. By making these dynamics experientially legible, the paper positions design fiction and legal exegesis as an exploratory method for examining how regulatory assumptions encounter AI in practice, inviting further work at the intersection of design and socio-legal research to address the challenges

of governing AI systems whose legal and risk profiles shift dynamically through situated interaction.

Unlike traditional legal commentary, which would focus on the text of the AI Act, or compliance auditing, which focuses on post-hoc outputs, this "fiction and exegesis" method allows for the interrogation of interface governance. It renders visible how the law might be operationalized into buttons, pauses, and prompts of a system in ways that vary from original intentions and assumptions.

We also see some potential in this method for serving as a "Pre-deployment Red Teaming" phase for organizations performing Fundamental Rights Impact Assessments. Future work should explore its application in other high-risk sectors like healthcare or education, where agentic systems similarly create regulatory ambiguity.

Acknowledgments

This publication was supported by the AI Futures Lab on Rights and Justice which is part of the TU Delft AI Labs program.

References

- [1] Ramón Alvarado. 2023. AI as an Epistemic Technology. *Science and Engineering Ethics* 29, 5 (Aug. 2023), 32. doi:10.1007/s11948-023-00451-3
- [2] James Auger. 2013. Speculative Design: Crafting the Speculation. *Digital Creativity* 24, 1 (March 2013), 11–35. doi:10.1080/14626268.2013.767276
- [3] John Bowers. 2012. The Logic of Annotated Portfolios: Communicating the Value of 'Research through Design'. In *Proceedings of the Designing Interactive Systems Conference (DIS '12)*. Association for Computing Machinery, New York, NY, USA, 68–77. doi:10.1145/2317956.2317968
- [4] Tara Brabazon, Jamie Quinton, and Narelle Hunter. 2022. The Scientist, the Artefact and the Exegesis: Challenging the parameters of the PhD. *International Journal of Creative and Arts Studies* 9, 1 (2022), 47–68.
- [5] Francesca Casnati, Alessandro Ianniello, and Alessia Romani. 2024. Provocation Through Narratives: New Speculative Design Tools for Human-Non-Human Collaborations. In *Multidisciplinary Aspects of Design*, Francesca Zanella, Giampiero Bosoni, Elisabetta Di Stefano, Gioia Laura Iannilli, Giovanni Matteucci, Rita Messori, and Raffaella Trocchianesi (Eds.). Springer Nature Switzerland, Cham, 747–755. doi:10.1007/978-3-031-49811-4_71
- [6] EunJeong Cheon and Norman Makoto Su. 2018. Futuristic Autobiographies: Weaving Participant Narratives to Elicit Values around Robots. In *Proceedings of the 2018 ACM/IEEE International Conference on Human-Robot Interaction (HRI '18)*. Association for Computing Machinery, New York, NY, USA, 388–397. doi:10.1145/3171221.3171244
- [7] European Commission. 2024. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act) (Text with EEA relevance). <http://data.europa.eu/eli/reg/2024/1689/oj/eng> Legislative Body: CONSIL, EP.
- [8] Mikkel Flyverbom. 2022. Overlit: Digital Architectures of Visibility. *Organization Theory* 3, 3 (July 2022), 26317877221090314. doi:10.1177/26317877221090314
- [9] Bill Gaver and John Bowers. 2012. Annotated Portfolios. *interactions* 19, 4 (July 2012), 40–49. doi:10.1145/2212877.2212889
- [10] Delaram Golpayegani, Harshvardhan J. Pandit, and Dave Lewis. 2023. To Be High-Risk, or Not To Be—Semantic Specifications and Implications of the AI Act's High-Risk AI Applications and Harmonised Standards. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency (Chicago, IL, USA) (FAccT '23)*. Association for Computing Machinery, New York, NY, USA, 905–915. doi:10.1145/3593013.3594050
- [11] Sandy J. J. Gould. 2024. Stochastic Machine Witnesses at Work: Today's Critiques of Taylorism are Inadequate for Workplace Surveillance Epistemologies of the Future. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems (CHI '24)*. Association for Computing Machinery, New York, NY, USA, 1–12. doi:10.1145/3613904.3642206
- [12] Hilmy Hanif, Jorge Constantino, Marie-Therese Sekwenz, Michel Van Eeten, Jolien Ubacht, Ben Wagner, and Yury Zhauniarovich. 2024. Navigating the EU AI Act Maze using a Decision-Tree Approach. *ACM Journal on Responsible Computing* 1, 3 (Sept. 2024), 1–16. doi:10.1145/3677174
- [13] Kyriakos Kyriakou and Jahna Otterbacher. 2023. In Humans, We Trust. *Discover Artificial Intelligence* 3, 1 (Dec. 2023), 44. doi:10.1007/s44163-023-00092-2
- [14] Ida Larsen-Ledet and Siân Lindley. 2022. Ways of seeing and being seen: People in the algorithmic knowledge base. In *Workshop at the 20th Eur. Conf. Comput.-Supported Cooperative Work*.

- [15] Q. Vera Liao and Michael Muller. 2019. Enabling Value Sensitive AI Systems through Participatory Design Fictions. arXiv:1912.07381 [cs] doi:10.48550/arXiv.1912.07381
- [16] Chaeyun Lim, Rabindra Ratan, Maxwell Foxman, David Beyea, David Jeong, and Alex P. Leith. 2025. Examining Attitudes about the Virtual Workplace: Associations between Zoom Fatigue, Impression Management, and Virtual Meeting Adoption Intent. *PLOS ONE* 20, 2 (Feb. 2025), e0312354. doi:10.1371/journal.pone.0312354
- [17] Jyoti Mann and Aaron Holmes. 2026. Meta Officially Ties Employee Performance to AI Usage; Microsoft on OpenClaw Security Risks. The Information. <https://www.theinformation.com/newsletters/applied-ai/meta-officially-ties-employee-performance-ai-usage-microsoft-openclaw-security-risks>
- [18] Mahan Mehrvarz, Dave Murray-Rust, Himanshu Verma, and Ben Wagner. 2025. Framing the (in)Visible: Insights into Visibility Practices of Remote Knowledge Workers. In *Proceedings of the 4th Annual Symposium on Human-Computer Interaction for Work (CHIWORK '25)*. Association for Computing Machinery, New York, NY, USA, 1–14. doi:10.1145/3729176.3729178
- [19] Microsoft. 2025. People Skills AI Inferencing. <https://learn.microsoft.com/en-us/copilot/microsoft-365/people-skills-ai-inferencing>.
- [20] George Mousourakis. 2014. Roman Law, Medieval Jurisprudence and the Rise of the European Ius Commune: Perspectives on the Origins of the Civil Law Tradition. *House of Riron* 47, 2 (2014), 22–84. https://www.researchgate.net/publication/367964750_Roman_Law_Medieval_Jurisprudence_and_the_Rise_of_the_European_Ius_Commune_Perspectives_on_the_Origins_of_the_Civil_Law_Tradition
- [21] Michael Muller, Justin D. Weisz, Stephanie Houde, and Steven I. Ross. 2024. Drinking Chai with Your (AI) Programming Partner: Value Tensions in the Tokenization of Future Human-AI Collaborative Work. In *Proceedings of the 3rd Annual Meeting of the Symposium on Human-Computer Interaction for Work (CHIWORK '24)*. Association for Computing Machinery, New York, NY, USA, 1–15. doi:10.1145/3663384.3663390
- [22] Xavier Parent-Rochelleau and Sharon K. Parker. 2022. Algorithms as Work Designers: How Algorithmic Management Influences the Design of Jobs. *Human Resource Management Review* 32, 3 (Sept. 2022), 100838. doi:10.1016/j.hrmr.2021.100838
- [23] Kat Roemmich, Florian Schaub, and Nazanin Andalibi. 2023. Emotion AI at Work: Implications for Workplace Surveillance, Emotional Labor, and Emotional Privacy. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems (CHI '23)*. Association for Computing Machinery, New York, NY, USA, 1–20. doi:10.1145/3544548.3580950
- [24] Zahara Rony and Heri Pardosi. 2021. Burnout digital monitoring on employee engagement at the company. *International Journal of Research in Business and Social Science* (2147- 4478) 10 (Nov. 2021), 156–162. doi:10.20525/ijrbs.v10i7.1412
- [25] Zoë Sadokierski. 2020. Developing Critical Documentation Practices for Design Researchers. *Design Studies* 69 (July 2020), 100940. doi:10.1016/j.destud.2020.03.002
- [26] Stefan Theil. 2025. Carefully Tailored: Doctrinal Methods and Empirical Contributions. *Oxford Journal of Legal Studies* 45, 4 (Dec. 2025), 1047–1075. doi:10.1093/ojls/gqaf029
- [27] Elvia Vasconcelos, Kristina Andersen, and Ron Wakkary. 2024. Who is holding the pen? Annotating collaborative processes. In *Proceedings of the 13th Nordic Conference on Human-Computer Interaction (Uppsala, Sweden) (NordiCHI '24)*. Association for Computing Machinery, New York, NY, USA, Article 65, 12 pages. doi:10.1145/3679318.3685401
- [28] Richmond Y. Wong and Vera Khovanskaya. 2018. Speculative Design in HCI: From Corporate Imaginations to Critical Orientations. In *New Directions in Third Wave Human-Computer Interaction: Volume 2 - Methodologies*, Michael Filimowicz and Veronika Tzankova (Eds.). Springer International Publishing, Cham, 175–202. doi:10.1007/978-3-319-73374-6_10
- [29] Richard Zhang and Duri Long. 2025. Beyond Content: Leaning on the Poetics of Defamiliarization in Design Fictions. *Proc. ACM Hum.-Comput. Interact.* 9, 1, Article GROUP5 (Jan. 2025), 19 pages. doi:10.1145/3701184