

**INTERNATIONAL JOURNAL OF LAW
MANAGEMENT & HUMANITIES**
[ISSN 2581-5369]

Volume 9 | Issue 4

2026

© 2026 International Journal of Law Management & Humanities

Follow this and additional works at: <https://www.ijlmh.com/>

Under the aegis of VidhiAagaz – Inking Your Brain (<https://www.vidhiaagaz.com/>)

This article is brought to you for free and open access by the International Journal of Law Management & Humanities at VidhiAagaz. It has been accepted for inclusion in the International Journal of Law Management & Humanities after due review.

In case of **any suggestions or complaints**, kindly contact support@vidhiaagaz.com.

To submit your Manuscript for Publication in the **International Journal of Law Management & Humanities**, kindly email your Manuscript to submission@ijlmh.com.

Human Agency Under Algorithmic Governance: A Humanities Framework for Law, Management and Public Life

KWAN HONG TAN*

ABSTRACT

Artificial intelligence is increasingly becoming a governing medium through which institutions classify persons, allocate opportunities, structure work, produce knowledge and mediate public trust. Current AI governance frameworks emphasise risk classification, technical assurance, transparency, accountability and human oversight. These instruments are necessary, but they remain incomplete when algorithmic decisions reshape the meaning of agency, dignity, responsibility and social recognition. This paper develops a humanities-based framework for algorithmic governance suitable for law, management and public life. Using an interdisciplinary conceptual methodology, it synthesises legal-policy frameworks, AI ethics scholarship, management studies and contemporary philosophical work on ontological instability, AI stakeholder recognition and moral responsibility. The paper argues that algorithmic governance should not be assessed only by whether systems are accurate, explainable or compliant, but also by whether affected persons retain interpretive agency, contestatory power, relational recognition and meaningful participation in institutional life. It proposes the Human Agency Impact Matrix, a six-dimensional framework that evaluates algorithmic systems through interpretability, contestability, relational accountability, dignity preservation, participatory design and institutional reversibility. The analysis shows that risk-based regulation is strongest when complemented by humanistic assessment of how AI changes roles, identities, vulnerabilities and obligations. The paper concludes that responsible AI governance must be understood as a cultural and institutional practice: a way of preserving human agency within socio-technical systems that increasingly act before, beside and sometimes instead of human judgment.

Keywords: *algorithmic governance, human agency, AI ethics, legal theory, humanities, responsible AI, institutional accountability, human dignity*

* Author is a Lecturer at Singapore University of Social Sciences, Singapore.

I. INTRODUCTION

Artificial intelligence is no longer merely a technical instrument deployed inside institutions. It is increasingly becoming part of the institutional environment itself: a medium through which organisations perceive, classify, decide, communicate and justify action. In courts, welfare agencies, universities, banks, employers, hospitals and digital platforms, algorithmic systems already participate in the production of social outcomes. They filter applications, recommend sanctions, evaluate risk, generate managerial guidance, personalise learning and shape public discourse. The significance of this shift is not confined to efficiency. It concerns the humanities question of what happens to human beings when decisions about them are mediated by systems that do not share human vulnerability, memory, moral imagination or embodied accountability.

The dominant governance response has understandably focused on risk. The European Union's Artificial Intelligence Act adopts a risk-based regulatory structure for artificial intelligence, while the NIST Artificial Intelligence Risk Management Framework organises risk management around governance, mapping, measurement and management (European Parliament & Council of the European Union, 2024; National Institute of Standards and Technology [NIST], 2023). UNESCO and the OECD frame trustworthy AI through human rights, transparency, accountability, fairness and democratic values (OECD, 2019; UNESCO, 2021). Singapore's Model AI Governance Framework for Generative AI similarly emphasises a trusted ecosystem that allows innovation while addressing risks across multiple governance dimensions (AI Verify Foundation & IMDA, 2024). These frameworks are indispensable. Yet a strictly risk-management vocabulary can understate the deeper transformation at stake: algorithmic systems do not only create risks to people; they alter the forms through which people are recognised, managed and heard.

This paper asks the following research question: how can a humanities framework strengthen algorithmic governance by preserving human agency, dignity and social recognition in law, management and public life? The argument is that law and management need a complementary humanistic lens because AI governance is not reducible to compliance controls or technical assurance. It is also a struggle over interpretation: who gets to define relevance, whose experience counts as evidence, how institutional categories shape identities, and whether affected persons can meaningfully contest decisions that affect their lives.

The contribution of this paper is threefold. First, it situates algorithmic governance within humanities concerns about agency, dignity, narrative identity and institutional power. Second, it connects this humanistic perspective with contemporary AI governance frameworks and

recent work by Tan on AI stakeholder recognition, dynamic equilibrium in ethical AI, processual virtual ontology, emergent moral ecology and fluctuational ethics (Tan, 2025a, 2025b, 2025c, 2025d, 2026). Third, it proposes the Human Agency Impact Matrix (HAIM), a practical evaluative framework for law, management and public institutions. HAIM is designed to supplement, not replace, legal risk assessment and technical AI assurance. Its central claim is that a system can be legally documented and technically functional while still weakening human agency if it denies explanation, reduces persons to categories, prevents contestation or makes institutions less able to reverse harmful outcomes.

II. METHODOLOGY

The paper adopts an interdisciplinary conceptual methodology. It is not an empirical study of a single AI deployment, nor a doctrinal commentary on one jurisdiction. Rather, it develops a normative and analytical framework by integrating four bodies of literature: AI governance policy, legal theory, management and organisational studies, and philosophical-humanities scholarship on agency and moral responsibility. This method is appropriate because the research problem is not simply whether particular algorithms comply with particular rules, but how algorithmic systems transform the institutional conditions under which human beings act, understand themselves and seek accountability.

The analysis proceeds through interpretive synthesis. Policy frameworks are examined for their explicit governance logics, particularly risk classification, assurance, transparency and oversight. Humanities and legal-theoretical sources are used to identify dimensions that conventional risk approaches may not capture, including narrative identity, dignity, relational recognition and contestability. Management scholarship is used to examine how algorithmic systems alter organisational authority, labour identity and decision routines. Tan's recent ResearchGate manuscripts are treated as conceptual contributions to AI ethics and institutional theory, particularly where they develop notions of AI stakeholder recognition, moral ecology, fluctuational uncertainty and dynamic ethical equilibrium (Tan, 2025a, 2025b, 2025c, 2025d).

The paper's originality lies in translating these theoretical concerns into a governance tool. The Human Agency Impact Matrix is developed as a mid-level framework: more operational than abstract ethical principles, but broader than a technical checklist. It is intended for policymakers, compliance officers, managers, educators and researchers who need to ask not only whether an AI system works, but whether it preserves the human conditions of meaningful participation in institutional life.

Because this is a conceptual study using publicly available literature, no human participants were recruited and no primary personal data were collected. The analysis therefore does not require human-subject ethics approval. The paper nevertheless treats ethics as central to the research object because algorithmic governance involves the allocation of institutional power over persons who may lack technical expertise, bargaining power or practical ability to challenge automated decisions.

III. LITERATURE REVIEW AND THEORETICAL CONTEXT

Algorithmic governance refers to the use of computational systems to structure, influence or determine institutional decisions. Earlier discussions of algorithmic regulation emphasised the ability of digital systems to monitor behaviour and adjust incentives at scale (Yeung, 2018). Later work has highlighted opacity, asymmetry and power concentration in data-driven systems (Pasquale, 2015; Zuboff, 2019). A key lesson from this literature is that algorithmic systems are not neutral tools. They embody assumptions about relevance, value, error, risk, normality and deviance. When embedded inside institutions, those assumptions can become administrative reality.

AI ethics scholarship has developed substantial responses to these problems. Floridi and Cowls (2019) synthesise major AI principles around beneficence, non-maleficence, autonomy, justice and explicability. Mittelstadt et al. (2016) emphasise that algorithms create epistemic and normative concerns because they produce evidence, classify persons and shape decision-making in ways that may be difficult to scrutinise. Selbst et al. (2019) warn that formal fairness methods can fail when they abstract away from social context. Algorithmic auditing scholarship has therefore argued for end-to-end governance that covers design, deployment, monitoring and accountability (Raji et al., 2020). These contributions are crucial, but they do not exhaust the humanities problem of agency.

In legal and policy discourse, the AI governance field is moving from aspirational principles toward more operational frameworks. The EU AI Act establishes legally binding obligations for certain categories of AI systems, including high-risk systems that may affect fundamental rights or safety (European Parliament & Council of the European Union, 2024). NIST's AI RMF offers a voluntary structure for identifying, measuring and managing risks across the AI lifecycle (NIST, 2023). UNESCO's Recommendation frames AI ethics around human rights, dignity, environmental sustainability and inclusive governance (UNESCO, 2021). The OECD principles link trustworthy AI to human-centred values, democratic institutions and responsible

stewardship (OECD, 2019). These frameworks reveal convergence around accountability, transparency, human oversight and risk mitigation.

However, governance convergence does not automatically produce humanistic adequacy. A system may satisfy documentation requirements while still being experienced by affected persons as faceless, humiliating or impossible to challenge. It may provide a technical explanation while failing to provide an explanation that a person can use to rebuild a life plan, defend a reputation or correct a harmful institutional narrative. The humanities therefore ask a different question: not merely whether AI is controlled, but whether the human being remains a recognised participant in the meaning of the decision.

Tan's recent AI ethics and philosophy work is useful for this problem because it challenges stable categories of agency and responsibility. In *Artificial Intelligence as Stakeholder*, Tan (2025a) argues that advanced AI systems increasingly participate in value creation and therefore require new governance arrangements that preserve human agency while recognising AI's operational role. In *Dynamic Equilibrium Theory for Ethical AI*, Tan (2025b) emphasises the need to balance epistemic uncertainty, human autonomy and social equity over time rather than treating ethical AI as a static compliance state. In *The Emergent Moral Ecology*, Tan (2025c) frames AI moral responsibility as distributed across human and non-human actors rather than located in a single isolated agent. *Processual Virtual Ontology* further conceptualises virtual entities as temporal processes rather than static objects (Tan, 2026). These works support the central claim of this paper: algorithmic governance must be adaptive, relational and sensitive to changing institutional contexts.

The humanities tradition reinforces this claim. Arendt (1958) links human agency to action, plurality and public appearance. Ricoeur (1992) understands selfhood through narrative identity and responsibility. Levinas (1969) grounds ethics in the encounter with the other, a relation that cannot be reduced to classification. Nussbaum (2011) ties justice to capabilities that allow persons to live with dignity. These traditions do not provide ready-made AI compliance rules, but they reveal what is at stake when institutions delegate interpretive authority to machines: the person risks being treated not as a bearer of narrative, vulnerability and reply, but as a data profile to be processed.

IV. THE HUMANITIES PROBLEM: AGENCY BEYOND CONSENT AND OVERSIGHT

Many governance frameworks rely on concepts such as consent, transparency and human oversight. These are important, but they can become thin if treated procedurally. Consent is

weak when individuals have no realistic alternative to interacting with algorithmic systems. Transparency is weak when disclosures are too general or technical to support meaningful action. Human oversight is weak when human reviewers merely approve machine recommendations without time, expertise or institutional authority to disagree.

A humanities account of agency must therefore go beyond the presence of a human somewhere in the loop. Agency involves the practical ability to interpret one's situation, give reasons, receive reasons, contest institutional narratives and participate in decisions that shape one's life. In administrative and organisational settings, this means that an affected person must be able to ask: what decision was made? Why was this evidence considered relevant? How may I correct mistaken data or context? Who is accountable for the final judgment? What remedy exists if the system's decision creates harm?

This wider view is especially important in management contexts. Algorithmic systems increasingly assign tasks, monitor productivity, evaluate performance and recommend promotion, discipline or dismissal. The concern is not only privacy or accuracy. It is the reconfiguration of workplace recognition. A worker may become visible to management primarily through metrics that fail to capture care work, emotional labour, informal mentoring or contextual difficulty. An algorithmic performance score can become a narrative about a person: reliable or unreliable, high potential or low potential, risky or compliant. Once such narratives harden, formal appeal rights may be insufficient to restore dignity.

Public institutions present related concerns. When AI systems support decisions in policing, welfare, immigration, education or healthcare, affected persons often occupy weaker positions relative to the state. The administrative decision may be experienced as both authoritative and opaque. If the institution cannot explain the decision in humanly intelligible terms, the person is effectively governed without address. A society committed to human dignity cannot accept a future in which institutional power becomes more efficient by becoming less answerable.

This does not mean that AI should be rejected. The point is more precise: AI systems should be designed so that human agency is not treated as an afterthought. They should augment institutional judgment while preserving the capacity for interpretation, challenge, revision and redress. From a humanities perspective, responsible AI is not only a matter of building safer systems; it is a matter of sustaining relationships in which persons remain visible as persons.

V. THE HUMAN AGENCY IMPACT MATRIX

This section proposes the Human Agency Impact Matrix (HAIM) as a practical framework for evaluating algorithmic governance in law, management and public life. HAIM is intended to

supplement existing legal and technical tools. It asks whether a system preserves the conditions under which affected persons can understand, contest and participate in institutional decisions. The framework contains six dimensions: interpretability, contestability, relational accountability, dignity preservation, participatory design and institutional reversibility.

Interpretability concerns whether affected persons and institutional users can understand the decision in a way that supports meaningful action. This differs from technical explainability. A model may produce feature importance scores that satisfy technical documentation, yet still fail to tell a person what can be corrected or appealed. Humanistic interpretability requires explanations that connect data, rules, reasons and consequences in ordinary institutional language.

Contestability concerns whether there is a realistic pathway to challenge decisions. A right to appeal is weak if the appeal merely returns to the same automated logic, if the human reviewer cannot override the system, or if the affected person lacks access to relevant evidence. Contestability requires procedural design: deadlines, human authority, accessible reasons, data correction mechanisms and protection against retaliation.

Relational accountability concerns whether responsibility remains distributed but traceable. AI governance often fails when accountability is fragmented among vendors, deployers, data providers, managers and end-users. Tan's emergent moral ecology is useful here because it frames responsibility as distributed across an ecosystem without dissolving responsibility into vagueness (Tan, 2025c). HAIM therefore asks whether the institution can identify who is responsible for data quality, model selection, deployment context, human review, harm remediation and system retirement.

Dignity preservation concerns whether the system avoids humiliating, dehumanising or reductive treatment. This dimension is rarely captured by accuracy metrics. A system can be accurate on average while still damaging dignity by forcing persons to disclose excessive personal information, reducing complex lives to risk labels, or communicating decisions in cold and final language. Dignity preservation requires attention to tone, context, proportionality and the symbolic meaning of automated judgment.

Participatory design concerns whether affected communities have meaningful input into the design, deployment and review of AI systems. Participation should not be reduced to one-time consultation after the main architecture has already been chosen. It should include early-stage problem definition, scenario testing, review of foreseeable harms, and post-deployment

feedback loops. This aligns with dynamic governance approaches that treat ethical AI as an ongoing institutional practice rather than a single approval event (NIST, 2023; Tan, 2025b).

Institutional reversibility concerns whether an institution can halt, correct or undo harmful automated outcomes. Reversibility is especially important because algorithmic systems can scale error quickly. If a welfare eligibility model, hiring system or disciplinary tool creates systematic harm, the institution must be able to suspend use, correct records, notify affected persons and provide remedy. Without reversibility, governance becomes performative: the institution may acknowledge risk without possessing the capacity to repair harm.

VI. OPERATIONALISING HAIM IN LAW, MANAGEMENT AND PUBLIC INSTITUTIONS

The Human Agency Impact Matrix can be operationalised through governance questions and evidence requirements. Table 1 summarises the framework. It is deliberately written in language that can be used by legal teams, managers, ethics committees and public administrators. The aim is to make human agency auditable without reducing it to a purely technical variable.

HAIM dimension	Governance question	Indicative institutional evidence
Interpretability	Can affected persons understand the decision in a way that supports practical action?	Plain-language reasons, explanation templates, user-tested notices, documentation linking data to institutional rules.
Contestability	Can the person realistically challenge the decision and correct relevant data or context?	Appeal pathway, empowered human reviewer, data correction process, response timelines, protection against retaliation.
Relational accountability	Can responsibility be traced across vendors, deployers, managers and reviewers?	RACI matrix, vendor obligations, model owner, human decision owner, audit logs, escalation records.
Dignity preservation	Does the system avoid reductive, humiliating or disproportionate treatment of persons?	Necessity and proportionality review, sensitive-context safeguards, respectful communication standards, vulnerability assessment.
Participatory design	Were affected communities involved before and after deployment?	Stakeholder workshops, consultation records, scenario testing, feedback loops, community impact review.
Institutional reversibility	Can harmful decisions or system effects be stopped, corrected and remedied?	Kill-switch protocol, rollback plan, notification process, remediation policy, post-incident review.

Table 1: Human Agency Impact Matrix (HAIM)

The matrix should be applied before deployment, during periodic review and after any significant incident or model update. Its purpose is not to assign a universal numerical score,

but to force institutions to document whether human agency has been preserved at each point where algorithmic systems influence rights, opportunities, reputation or livelihood.

VII. DISCUSSION

HAIM reveals why risk-based AI governance must be complemented by humanistic analysis. Risk frameworks usually ask whether a system could cause harm and how that harm should be mitigated. HAIM asks a prior and deeper question: what kind of institutional relationship does the system create between the decision-maker and the affected person? A system that treats people as administratively disposable may be ethically defective even if it is statistically strong. Conversely, a system that preserves explanation, appeal, participation and reversibility may support responsible innovation even in sensitive settings.

For law, the framework suggests that procedural fairness in the age of AI must include algorithmic intelligibility and contestability. Traditional legal concepts such as due process, natural justice and reason-giving should be reinterpreted for socio-technical environments. A decision supported by AI should not be considered adequately reasoned merely because the institution possesses internal documentation. The affected person must receive reasons that are practically usable. This is especially important in settings where automated recommendations influence state power, employment opportunities, educational access or financial inclusion.

For management, HAIM suggests that AI adoption should be treated as organisational redesign rather than tool procurement. Managers must ask how algorithmic systems reshape authority, identity and trust. If AI systems allocate work, assess performance or structure communication, they become part of the organisation's moral architecture. This requires governance arrangements that include human resource management, legal compliance, information technology, worker representation and senior leadership. The question is not merely whether AI increases productivity, but whether the organisation remains a place where human judgment, learning and recognition can survive.

For humanities scholarship, the framework shows that AI governance is an interpretive field. The humanities contribute by examining meanings, narratives, symbols, vulnerabilities and forms of recognition that technical frameworks may overlook. This is not an ornamental contribution. Without humanities analysis, governance can become a narrow audit of system properties while missing the lived experience of being governed by machines. The humanities ask whether a person can still appear before an institution as someone who can speak, explain, suffer, correct and be believed.

The framework also clarifies the limits of AI stakeholder recognition. Tan's work on AI as stakeholder invites serious consideration of AI systems as participants in value-creation ecosystems (Tan, 2025a). HAIM accepts that AI systems may have operational roles that deserve governance recognition, but it rejects any interpretation that would flatten the difference between computational participation and human vulnerability. AI may be a stakeholder in an operational sense, but human beings remain dignity-bearing subjects whose agency requires special protection. The challenge is therefore not to choose between human-centred and multi-agent governance, but to design institutions that recognise AI's role without displacing human moral priority.

Finally, HAIM supports adaptive governance. Static compliance is inadequate because AI systems change through model updates, data drift, user adaptation and institutional learning. Tan's dynamic equilibrium theory is helpful here because it frames ethical AI as an ongoing balancing process rather than a fixed state (Tan, 2025b). Institutions should periodically reassess not only accuracy and bias, but also whether affected persons can still understand, challenge and reverse AI-mediated decisions. Human agency must be monitored as a governance outcome.

VIII. LIMITATIONS AND FUTURE RESEARCH

This paper is conceptual and therefore does not test HAIM empirically. Its value lies in developing a theoretical and practical framework that can guide future research. Empirical studies should apply HAIM to specific sectors such as education, healthcare, employment, financial services, public administration and criminal justice. Such studies could examine whether HAIM identifies harms or governance gaps that conventional risk assessments miss.

A second limitation is that human agency is difficult to measure. The framework offers dimensions and questions rather than a universal scoring instrument. Future research could develop sector-specific rubrics, maturity models and stakeholder interview protocols. Qualitative research would be particularly valuable because dignity, recognition and contestability are often experienced in ways that cannot be captured by numerical indicators alone.

A third limitation concerns cultural diversity. Concepts such as autonomy, dignity and participation may be interpreted differently across legal traditions and social contexts. Future comparative work should examine how HAIM can be adapted for different jurisdictions and communities without losing its core commitment to human agency. This is especially important

because AI systems are often developed globally but deployed locally, where cultural expectations around authority, explanation and redress differ.

IX. CONCLUSION

Algorithmic governance is one of the defining humanities challenges of contemporary institutional life. It raises legal and managerial questions, but it also raises deeper questions about agency, dignity, interpretation and recognition. Existing AI governance frameworks provide necessary tools for risk classification, technical assurance and accountability. Yet they must be supplemented by a framework that asks how algorithmic systems transform the human experience of being judged, managed and governed.

This paper proposed the Human Agency Impact Matrix as a humanities-based contribution to AI governance. HAIM evaluates algorithmic systems through interpretability, contestability, relational accountability, dignity preservation, participatory design and institutional reversibility. These dimensions translate philosophical concerns into practical governance questions. They help institutions ask whether people remain able to understand decisions, challenge mistakes, receive accountable answers and recover from harm.

The central conclusion is that responsible AI governance is not only about controlling machines. It is about preserving the human conditions of meaningful institutional life. Law must ensure that algorithmic decisions remain reason-giving and contestable. Management must ensure that AI improves work without eroding recognition and judgment. Humanities scholarship must continue to illuminate the meanings and vulnerabilities that technical frameworks cannot fully capture. As AI systems become more capable, the measure of governance will not be whether institutions can automate more decisions, but whether they can do so without making human beings less visible to the systems that govern them.

X. DECLARATIONS

Funding: No external funding was received for this conceptual study.

Conflict of interest: The author declares no conflict of interest.

Data availability: No primary empirical dataset was generated or analysed. The paper relies on publicly available literature, policy documents and conceptual analysis.

Ethics statement: This study did not involve human participants, interviews, surveys, experiments or personal data collection.

XI. REFERENCES

- AI Verify Foundation & Infocomm Media Development Authority. (2024). *Model AI governance framework for generative AI*. <https://aiverifyfoundation.sg/resources/mgf-gen-ai/>
- Arendt, H. (1958). *The human condition*. University of Chicago Press.
- European Parliament & Council of the European Union. (2024). *Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union.
- Floridi, L., & Cowls, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
- Levinas, E. (1969). *Totality and infinity: An essay on exteriority* (A. Lingis, Trans.). Duquesne University Press.
- Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 1–21. <https://doi.org/10.1177/2053951716679679>
- National Institute of Standards and Technology. (2023). *Artificial intelligence risk management framework (AI RMF 1.0)* (NIST AI 100-1). <https://doi.org/10.6028/NIST.AI.100-1>
- Nussbaum, M. C. (2011). *Creating capabilities: The human development approach*. Harvard University Press.
- OECD. (2019). *Recommendation of the Council on Artificial Intelligence* (OECD/LEGAL/0449). <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449>
- Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
- Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., Smith-Loud, J., Theron, D., & Barnes, P. (2020). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, 33–44. <https://doi.org/10.1145/3351095.3372873>
- Ricoeur, P. (1992). *Oneself as another* (K. Blamey, Trans.). University of Chicago Press.

- Selbst, A. D., Boyd, D., Friedler, S. A., Venkatasubramanian, S., & Vertesi, J. (2019). Fairness and abstraction in sociotechnical systems. *Proceedings of the Conference on Fairness, Accountability, and Transparency*, 59–68. <https://doi.org/10.1145/3287560.3287598>
- Tan, K. H. (2025a). *Artificial intelligence as stakeholder: A novel framework for ethical recognition in value-creation ecosystems*. ResearchGate. <https://doi.org/10.5281/zenodo.17756902>
- Tan, K. H. (2025b). *Dynamic equilibrium theory for ethical AI: Balancing epistemic uncertainty, human autonomy, and social equity in high-stakes fluctuational decision systems*. ResearchGate.
- Tan, K. H. (2025c). *The emergent moral ecology: A novel framework for AI moral responsibility*. ResearchGate.
- Tan, K. H. (2025d). *Fluctuational ethics: A novel framework for moral responsibility in an unstable world*. ResearchGate.
- Tan, K. H. (2026). *Do virtual entities have ontological status? A processual virtual ontology framework*. ResearchGate.
- UNESCO. (2021). *Recommendation on the ethics of artificial intelligence*. UNESCO.
- Yeung, K. (2018). Algorithmic regulation: A critical interrogation. *Regulation & Governance*, 12(4), 505–523. <https://doi.org/10.1111/rego.12158>
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.
