

Balancing Innovation and Harmony: A Daoist Approach to Artificial Intelligence Law

Ramandeep Kaur¹, Dr. Reena Bishnoi², Dr. Dhruv³, Dr. Jogiram Sharma⁴, Mr. Kapil Dev⁵, Dr. Oindrila Bhattacharya⁶

¹Assistant Professor of Law

Centre for Legal Studies, Gitarattan International Business School, Rohini, Delhi

Email: d.raman17@yahoo.com

²Professor (Law)

Swami Vivekanand Subharti University, Meerut

Email: bishnoireena@gmail.com

³Assistant Professor

Dr. Akhilesh Das Gupta Institute of Professional Studies, GGSIPU

Email: llbdhruvadvacate@gmail.com

⁴Professor

Dr. Akhilesh Das Gupta Institute of Professional Studies, GGSIPU

Email: jrsharma1964@gmail.com

⁵Assistant Professor

J. C. College of Law, MDU

Email: kapillawlearner2026@gmail.com | Mobile: 9253388262

⁶Assistant Professor,

Department of English & Literary Studies Brainware University

Email: oindrila.bhattacharya9@gmail.com

Abstract

Contemporary artificial intelligence (AI) governance is dominated by two regulatory postures: the European Union's precautionary, risk-tiered rulebook and the more permissive, innovation-first posture associated with the United States. Both largely treat law as an external wall built to contain a technology assumed to be inherently unruly. This paper asks what classical Daoist philosophy, and in particular the concepts of wu wei (non-coercive action), ziran (spontaneous naturalness), and the dynamic complementarity of yin and yang, might offer as a corrective to that assumption. Drawing on the Daodejing, the Zhuangzi, and recent scholarship on Daoism and technology, the paper develops a reading of regulation as a channel that guides a technology's own tendencies rather than a barrier erected against them. It then examines China's AI governance regime as a partial and genuinely contested illustration of harmony-oriented rhetoric operating alongside far less Daoist mechanisms of state control, before turning to the central objections against importing a two-thousand-year-old philosophical tradition into contemporary legal design: the risk of romanticizing an essentialized 'East,' and the risk that a language of harmony can be made to justify opacity rather than restraint. The paper concludes that Daoist thought is most useful not as a ready-made regulatory template but as a disposition -- a standing reminder that adaptive humility, rather than exhaustive control, may be the more durable foundation for governing a technology whose trajectory no single actor, East or West, can fully foresee.

Keywords: Daoism; wu wei; artificial intelligence law; AI governance; regulatory harmony; comparative philosophy of technology

1. Introduction

There is a peculiar sameness to how the world's major jurisdictions have chosen to talk about artificial intelligence, even where their substantive rules diverge sharply. The vocabulary is one of containment: risk tiers, red lines, guardrails, prohibited practices, kill switches. The European Union's Artificial Intelligence Act, the first comprehensive horizontal AI statute in the world, sorts systems into unacceptable, high, limited,

and minimal risk categories and attaches escalating obligations to each (European Parliament and Council of the European Union, 2024; White & Case, 2024). The implicit picture is of a technology that behaves, left to itself, somewhat like flood water: left unchecked it will find every crack, so the law's task is to build walls of the right height in the right places. It is a coherent picture, and in many respects a defensible one. But it is also, this paper suggests, only one picture, and importing it uncritically as the picture of what AI regulation must look like forecloses other ways of thinking about the relationship between a fast-moving technology and the law that is meant to govern it.

Classical Daoism offers a genuinely different starting point. Where the dominant regulatory grammar treats order as something imposed on a recalcitrant world, the *Daodejing* and the *Zhuangzi* treat order as something that already exists, latent, in the way things naturally unfold, and treat the ruler's or the sage's task as one of alignment rather than construction (Lao Tzu, 1963; Graham, 2001). The Daoist term for this alignment, *wu wei*, is routinely mistranslated as passivity; it is closer to acting in a way so well fitted to the grain of a situation that no visible strain is required to produce the result (Musacchio, 2025). This paper asks what it would mean to take that idea seriously as a resource for AI law, not as an exotic ornament attached to a Western regulatory skeleton, but as a genuinely different account of what regulation is for.

The argument proceeds in five steps. Section 2 lays out the relevant Daoist concepts -- the Dao itself, *wu wei*, *ziran*, and the yin-yang logic of complementary opposites -- together with the classical tradition's own, considerably more ambivalent, engagement with technology, most vividly captured in the *Zhuangzi*'s story of the gardener who refuses a labor-saving machine for fear of what it will do to his heart (Krajewska, n.d.; Wheeler, 2022). Section 3 sketches the two regulatory grammars that presently dominate global AI governance discourse, the European Union's precautionary model and the more permissive posture associated with the United States, both of which this paper reads as variants of the same control-oriented assumption about what law is for. Section 4 brings the two together, asking where Daoist categories illuminate blind spots in the control paradigm and where, just as importantly, they run up against the genuine hazards -- discrimination, mass surveillance, algorithmic harm -- that the control paradigm exists to address. Section 5 turns to China, whose official AI governance documents borrow the vocabulary of harmony and human-machine coexistence while operating a regulatory apparatus that is, on close inspection, considerably more directive and state-centered than the philosophical language would suggest (Beijing Academy of Artificial Intelligence, 2019; State Council of the People's Republic of China, 2017). Section 6 confronts the most serious objections to the whole project, and Section 7 proposes, more modestly, a set of dispositions rather than rules that a Daoist sensibility might contribute to AI law going forward.

2. The Daoist Conceptual Toolkit

Daoism is not a single doctrine so much as a family of related texts and practices that took shape in China from roughly the fourth century BCE onward, of which the two foundational works are the *Daodejing*, traditionally attributed to Laozi, and the *Zhuangzi*, a looser and more literary collection associated with the philosopher of the same name (Coutinho, 2014). Both texts orbit a term, Dao (道), that resists tidy translation -- 'the Way' is the conventional gloss, but the word carries a sense of process and unfolding that 'way' alone does not capture. The Dao is not a rule to be followed or a lawgiver to be obeyed; it is closer to the pattern according to which things already tend to happen when nothing artificial interferes with them (Coutinho, 2014).

From this cosmology follows the concept most often invoked in discussions of Daoism and governance: *wu wei*, literally 'non-action,' though the classical texts make clear that literal inaction is not what is meant. Laozi's famous formulation is that the Dao 'never acts, yet nothing is left undone' (Lao Tzu, 1963, ch. 37), a paradox that only resolves once *wu wei* is understood as action so well matched to circumstance that it requires no visible exertion and generates no resistance. Musacchio (2025) traces the political application of this idea to the early Han dynasty's *huang-lao* statecraft, a period in which rulers deliberately practiced light taxation and minimal bureaucratic interference on the theory that a well-ordered society, like a healthy body, tends toward equilibrium if the ruler does not constantly disturb it. The companion concept is *ziran* (自然), usually rendered 'naturalness' or 'spontaneity,' which names the state of a thing developing according to its own internal tendency rather than an externally imposed design (Coutinho, 2014). This emphasis on unforced naturalness as the seat of virtue is also what most sharply distinguishes Daoist ethics from the rival Confucian tradition, which locates virtue instead in cultivated ritual propriety and structured social role -- a contrast worth flagging because Confucian, not Daoist, harmony is more often what gets invoked, sometimes loosely, in official Chinese governance rhetoric, a distinction taken up in Section 5 (College of Western Idaho, n.d.). A third structuring idea, less a discrete

term than a way of seeing, is the yin-yang logic of complementary opposition: innovation and restraint, speed and deliberation, openness and boundary are not, on this view, a war to be won by one side, but a pairing whose health depends on each pole checking and completing the other (Zhu, n.d.).

It would be a mistake, however, to read classical Daoism as straightforwardly pro-technology simply because it distrusts heavy-handed control. The tradition's most cited episode on the question runs in the opposite direction. In a passage from the Zhuangzi's 'Heaven and Earth' chapter, Confucius's disciple Zigong comes across an old gardener who is laboriously hauling water by hand and suggests he adopt a well-sweep, a simple lever device that would let him irrigate far more ground for far less effort. The gardener, Watson's (2013) translation records, flushes with something between anger and amusement and refuses: 'where there are machines, there are bound to be machine worries; where there are machine worries, there are bound to be machine hearts. With a machine heart in your breast, you've spoiled what was pure and simple ... I am not ignorant of this device -- I would be ashamed to use it' (Watson, 2013, p. 91; see also Krajewska, n.d.). The gardener's worry is not that the machine will malfunction or be misused; it is that adopting a tool reorganizes the user's own mind around the tool's logic, and that this reorganization, once underway, is very hard to reverse (Wheeler, 2022). Read this way, the parable is less a Luddite manifesto than an early and remarkably durable statement of a concern that now animates debates over algorithmic dependency and the erosion of independent judgment: that the danger of a technology may lie less in its failures than in its successes, in how thoroughly it comes to reshape the mind that relies on it (Allen, 2010; Wheeler, 2022).

The relevance of this parable to AI law is not that it counsels rejecting AI, any more than the gardener's refusal represents the settled Daoist position on technology as such -- the Zhuangzi elsewhere celebrates figures like the butcher Cook Ding, whose mastery of his knife is itself a form of *wu wei*, technique so thoroughly internalized that it no longer feels like technique (Graham, 2001; Wheeler, 2022). The relevance, rather, is diagnostic: the tradition asks a question that contemporary AI law asks only occasionally and unsystematically, namely what a given technology does to the disposition of the people and institutions that come to depend on it, independent of whether any particular use of it causes a discrete, attributable harm. A regulatory grammar built entirely around discrete, attributable harms -- which is broadly what the risk-tier model of the EU AI Act does -- may simply lack the vocabulary to ask that question at all (European Parliament and Council of the European Union, 2024).

3. Two Dominant Grammars of AI Law

The European Union's Artificial Intelligence Act, Regulation (EU) 2024/1689, entered into force on 1 August 2024 and is being phased in through 2027 (White & Case, 2024). It sorts AI systems into four tiers -- unacceptable risk, which is banned outright; high risk, which triggers a dense set of obligations around data governance, human oversight, and post-market monitoring; limited risk, which triggers narrower transparency duties; and minimal risk, which is essentially left alone -- and backs the scheme with fines of up to seven percent of global turnover (European Parliament and Council of the European Union, 2024). The underlying theory of regulation is precautionary and classificatory: identify in advance the categories of harm that matter, sort systems into the category whose harm-potential they most resemble, and calibrate the intensity of legal intervention accordingly. It is, in its own terms, a considered and often admirable structure. It is also unmistakably a wall-building exercise: a taxonomy of dangers, matched to a taxonomy of controls, imposed from outside onto a technology treated as inherently in need of restraint (Yeung, Howes, & Pogrebnia, 2020).

The United States has so far declined to enact anything resembling a horizontal AI statute, relying instead on sector-specific rules, agency guidance, and, since the change in federal administration, an explicit strategic emphasis on maintaining AI leadership through comparatively light federal intervention (Wu & Liu, 2023). This looks, superficially, like the opposite of the European model, and in terms of the volume and density of binding rules it certainly is. But the underlying grammar is arguably continuous with the European one rather than a genuine alternative to it: both treat regulation as a lever to be set at a particular strength -- heavier in the EU, lighter in the US -- along a single dial running from restraint to permission. Neither asks whether the dial itself, restraint versus permission, is the right axis along which to think about the problem. Global surveys of AI ethics guidelines document the same pattern at the level of stated principles: an overwhelming convergence, across dozens of national and corporate frameworks, on a similar list of values -- transparency, fairness, accountability, non-maleficence -- with comparatively little variation in what those values are taken to mean or in the underlying theory of why law is the right instrument to secure them (Jobin, Ienca, & Vayena, 2019; Floridi et al., 2018; Floridi & Cowls, 2019).

This is not a criticism of either jurisdiction's substantive choices, several of which respond to genuine

and well-documented harms. It is an observation about grammar: both models proceed from the assumption that a technology, left alone, tends toward disorder, and that the state's task is to impose an external architecture of permission and prohibition strong enough to correct for that tendency. UNESCO's Recommendation on the Ethics of Artificial Intelligence, adopted by acclamation across a genuinely global membership in 2021, gestures toward a broader set of values, including human oversight, sustainability, and the flourishing of diverse cultural and epistemic traditions, but even this document is structured as a set of external commitments states are asked to adopt, rather than as an inquiry into whether the imposition of external commitments is, by itself, an adequate theory of technology governance (UNESCO, 2021).

4. Where Daoism Meets AI Law: Convergences and Tensions

Read against this backdrop, *wu wei* offers something more specific than a general plea for lighter regulation. Its political application in Han-dynasty statecraft was not an absence of governance but a particular theory of how governance works best: through calibrated, minimal, well-timed intervention that works with the grain of existing social and economic tendencies rather than against them, on the wager that heavy intervention tends to generate exactly the instability it is meant to prevent (Musacchio, 2025). Translated into AI law, this suggests a test that the risk-tier model does not naturally ask: not simply how severe is the potential harm, but how well does a given rule fit the actual technical and organizational tendencies of the systems it targets, and how much of the rule's apparent stringency is, on inspection, adding friction without adding protection. Regulatory sandboxes, of the kind the EU AI Act itself provides for, are one institutional expression of something like this logic -- a space in which oversight adapts to the technology's own developmental path rather than being fixed in advance (European Parliament and Council of the European Union, 2024).

Ziran, the principle of spontaneous, self-so development, cuts in a similar but distinct direction, toward what technologists have started to call 'wu wei-effectiveness' or biomimetic design: building systems and rules that work with a technology's actual affordances rather than forcing it into an artificial mold (Wheeler, 2022, drawing on Allen, 2010). Applied to AI governance, this might counsel favoring standards that specify outcomes -- a model must not discriminate on protected characteristics, a system must be explicable to the degree its stakes require -- over standards that specify particular architectures or processes, on the theory that architecture-specific rules calcify quickly in a field whose technical substrate changes faster than any legislature can track. It is a case, in other words, for a certain humility about the regulator's own capacity to anticipate the shape technology will take, and for building that humility into the rule's design rather than treating it as a temporary embarrassment to be corrected by ever more detailed annexes.

The yin-yang frame is, of the three, the most directly applicable to the innovation-versus-restraint debate that structures so much AI policy discourse, precisely because it refuses the framing of that debate as a debate. Innovation and restraint are not, on this view, opposing teams competing for a larger share of the regulatory dial; they are complementary functions, each of which degrades without the other. A regime that maximizes only innovation is not more free of restraint in any meaningful sense -- it simply relocates the restraint downstream, onto whoever absorbs the resulting harm, and eventually onto the industry itself once public trust collapses. A regime that maximizes only restraint is not more safe in any meaningful sense either, since excessive friction pushes development into less visible, less accountable jurisdictions and actors (Wu, 2023). The Daoist contribution here is not a new balance point on the existing dial but a reframing of what balance is for: not a compromise between two opposed goods, but the recognition that innovation and harm-prevention are, properly understood, two aspects of the same underlying process, and that a rule which damages one will, sooner or later, damage the other.

None of this is to suggest that Daoist categories map neatly onto every problem AI law must solve. The gardener's worry about the 'machine heart' is a worry about cultivation and disposition; it has comparatively little to say about the concrete, distributable harms -- discriminatory lending algorithms, opaque hiring tools, mass biometric surveillance -- that motivate the EU AI Act's high-risk category and that a purely dispositional ethic of non-interference is not obviously equipped to prevent (Yeung, Howes, & Pogrebná, 2020). Cross-cultural scholarship on AI ethics has been explicit about this limitation: a governance framework built entirely on cultivated virtue and self-restraint, without any binding external check, has historically proven fragile precisely where the incentives to cut corners are strongest (ÓhÉigeartaigh, Whittlestone, Liu, Zeng, & Liu, 2020). A Daoist-informed approach to AI law, then, is best read not as a replacement for rules against discrete, attributable harm, but as a supplementary disposition that governs how those rules are designed, how heavily they are enforced, and how much residual discretion is left for a system, and the people who build and use it, to find their own fit with the world -- a disposition, in short,

about the manner of regulating, layered onto rather than substituted for the substance of what is regulated.

5. Harmony as Governance Rhetoric: The Case of China

China offers the most visible contemporary attempt to fold classical Chinese philosophical vocabulary, including Daoist and Confucian registers of harmony, into official AI governance discourse, and for that reason it is worth examining closely, and skeptically. The State Council's New Generation Artificial Intelligence Development Plan of 2017 set out the ambition of global AI leadership by 2030 and explicitly framed AI development as something to be pursued in a manner compatible with social stability and human welfare, language that resonates with, without directly quoting, the classical ideal of harmony between human effort and natural order (State Council of the People's Republic of China, 2017; China Briefing, 2023). The Beijing AI Principles, released in 2019 by a coalition of leading Chinese universities and technology firms, went further, framing the core aim of AI development as 'beneficial AI for humankind and nature' and calling for human-AI coordination rather than a purely competitive or extractive relationship between the two (Beijing Academy of Artificial Intelligence, 2019). Diplo's (2025) account of a recent Beijing dialogue on AI governance and philosophy captures this register directly, describing *wu wei*, reinterpreted for a contemporary audience, as guidance toward 'non-coercive action or effortless alignment with the natural flow of things' rather than passivity or fatalism.

Set against this vocabulary of harmony, however, is a regulatory apparatus that is considerably more directive than the philosophical framing suggests. China's Interim Measures for the Management of Generative Artificial Intelligence Services, jointly issued by seven ministries and in effect since August 2023, require providers of public-facing generative AI to undergo security assessments, register their algorithms with the Cyberspace Administration of China, ensure that training data and outputs align with 'core socialist values,' and cooperate with government inspection on demand (Cyberspace Administration of China et al., 2023; Ashurst, 2023; China Briefing, 2023). This is not obviously *wu wei* in the Han-dynasty, light-touch sense; it is closer to an assertively steered model in which the state retains extensive, continuously exercisable authority over what AI systems may say and do, layered underneath a discourse of natural harmony and beneficial coexistence. S. S. Wu (2023) reads this combination -- expansive developmental ambition, an official rhetoric of balance and human-machine harmony, and a dense, centrally coordinated apparatus of content and data control -- as characteristic of Beijing's broader approach to technology governance, in which philosophical language legitimates rather than constrains state authority.

The lesson for this paper's argument is a cautionary one. Invoking Daoist vocabulary does not, by itself, produce a Daoist regulatory outcome, any more than invoking the language of liberty guarantees a liberal one. Harmony can describe a genuine disposition of restraint, of the kind the Han-dynasty *huang-lao* rulers are credited with practicing (Musacchio, 2025), or it can describe an aspiration toward social conformity imposed from above, in which 'harmony' functions less as a limit on state power than as a justification for its exercise. Distinguishing these two uses requires looking past the vocabulary to the actual distribution of discretion -- who retains the power to intervene, on what triggers, and subject to what check -- a distinction the philosophical texts themselves, incidentally, were alert to: the *Daodejing* is at least as preoccupied with warning rulers against the appearance of virtue masking domination as it is with commending harmony as such (Lao Tzu, 1963).

6. Objections and Cautions

Three objections deserve direct engagement rather than a passing acknowledgment. The first is the charge of orientalism: that treating 'Daoism' as a coherent, exportable alternative to 'Western' regulatory thinking flattens both a philosophical tradition spanning over two millennia of internal debate and a regulatory landscape that is itself far from monolithic, and that doing so risks reproducing an old and unflattering habit of using a romanticized East as a foil for critiquing the West. This risk is real, and the honest response is not to deny it but to narrow the claim. This paper does not argue that China, or Chinese civilization, embodies Daoist principles in its governance practice -- Section 5 argues close to the opposite. Nor does it argue that *wu wei* is untranslatable into, or absent from, other traditions; light-touch, adaptive regulatory philosophies have plenty of Western analogues, from common-law incrementalism to regulatory-sandbox experimentation. The claim is narrower and, the paper hopes, more defensible: that a specific, textually grounded set of concepts from the Daoist canon names a disposition toward technology governance that the currently dominant AI-law vocabulary does not name well, and that naming it is useful regardless of the geographic label attached to it.

The second objection is that a philosophy built around minimal intervention is simply the wrong tool for a technology capable of large-scale, fast-moving, and difficult-to-reverse harm. Facial recognition

deployed at national scale, algorithmic decision systems embedded in credit, employment, and criminal justice, and generative systems capable of large-scale disinformation are not, on this view, problems that adaptive humility is equipped to solve; they are problems that require binding, enforceable, externally imposed limits precisely because the actors who build and deploy such systems cannot be relied upon to find the harmonious fit themselves (Yeung, Howes, & Pogrebnia, 2020). This objection is largely correct, and the response offered here is not to contest it but to relocate the argument: nothing in this paper's reading of wu wei or ziran counsels against binding prohibitions on the clearest and most severe harms, of the kind the EU AI Act's unacceptable-risk category already targets. The Daoist contribution is to the much larger, harder-to-classify middle ground of AI systems whose risk profile is genuinely uncertain, where the choice is not between control and no control but between a dense, front-loaded, architecture-specific rulebook and a lighter, outcome-focused, adaptively monitored one.

The third objection concerns translation and fidelity. Wu wei and ziran are technical terms embedded in a specific metaphysics, one in which the Dao is understood as a real, cosmologically prior pattern that human effort can align with or obstruct; stripping the terms of that metaphysics and redeploying them as generic synonyms for 'light-touch regulation' arguably does violence to the tradition even as it borrows its vocabulary. This is a fair charge, and this paper does not claim to resolve it, only to flag it honestly: what is offered here is a philosophically informed analogy rather than a claim that contemporary AI regulators should adopt classical Daoist cosmology as a governing premise. The value of the analogy lies in the questions it prompts a regulator to ask -- does this rule fit the technology's actual tendencies, or fight them; does this restriction address a real harm, or merely the appearance of control; is the balance between innovation and restraint being treated as a genuine complementarity or a zero-sum contest -- rather than in any claim that the Dao itself underwrites a specific policy conclusion.

7. Toward a Daoist-Informed Regulatory Disposition

Taken together, the preceding sections point toward a modest set of dispositions rather than a rival rulebook. First, a bias toward outcome-based rather than architecture-based rules wherever the harm in question can be specified with reasonable precision, on the ziran-inflected theory that rules calcifying around a particular technical architecture will misfit the technology faster than legislatures can amend them (Wheeler, 2022). Second, a preference for calibrated, reversible, and continuously monitored intervention -- sandboxes, staged rollouts, sunset clauses -- over front-loaded, exhaustively specified regimes, consistent with the wu wei intuition that governance works best when it can adjust to a technology's actual developmental path rather than a projected one fixed years in advance (Musacchio, 2025; European Parliament and Council of the European Union, 2024). Third, an explicit refusal to treat innovation and restraint as opposing values to be traded off against one another, in favor of treating them, in the yin-yang sense, as functions that each degrade without the other, which in practical terms means building safety review and innovation support into the same institutional process rather than assigning them to competing agencies with competing mandates (Zhu, n.d.). Fourth, and perhaps most distinctively, an attentiveness to dispositional as well as distributional harm -- to what a given AI system does to the judgment, autonomy, and skill of the people who come to depend on it, independent of whether any discrete decision it makes can be shown to be wrong -- a category of concern the Zhuangzi's gardener names precisely and that current AI law names only obliquely, if at all, through concepts like automation bias and deskilling (Krajewska, n.d.; Wheeler, 2022).

None of these dispositions is unique to Daoism, and the paper does not claim otherwise. What the Daoist frame adds is a coherent underlying rationale for holding them together: not four separate policy preferences but a single, textually grounded account of what good governance of a fast-moving, only partly predictable phenomenon looks like, developed originally for a very different kind of unruly complexity -- an empire, a river system, a human life -- and offered here, with appropriate caution about the distance being traveled, as one resource among several for thinking about a new one.

8. Conclusion

AI law, as it currently exists, speaks almost entirely in the grammar of containment: risk tiers, prohibited practices, mandatory disclosures, escalating fines. That grammar answers real and urgent problems, and nothing in this paper argues for its abandonment. But a technology whose trajectory no regulator, and arguably no developer, can fully foresee may be better served by a regulatory philosophy that expects to keep adjusting rather than one that expects, at the moment of enactment, to have gotten the architecture right. Classical Daoism, read carefully and without the romanticism that so often attaches to appeals to 'Eastern wisdom,' offers a vocabulary for that kind of adjustment: wu wei as calibrated, minimal,

well-timed intervention rather than passivity; ziran as a preference for rules that work with a technology's own tendencies rather than against them; yin-yang as a refusal to treat innovation and restraint as a zero-sum contest; and, from the Zhuangzi's gardener, a reminder that the deepest costs of a technology may be dispositional long before they are visible in any dataset. China's own experience shows how easily this vocabulary can be borrowed without being practiced, which is itself an argument for taking the underlying disposition seriously rather than treating the label as sufficient. The Dao that can be legislated, one is tempted to say, adapting Laozi's famous opening line, is not the eternal Dao -- but the impulse behind the line, that the wisest governance is the governance most alive to its own limits, may be exactly what a field still inventing its regulatory vocabulary needs to hear.

References

- Allen, B. (2010). A Dao of technology? *Dao: A Journal of Comparative Philosophy*, 9(1), 151-160. <https://doi.org/10.1007/s11712-010-9155-y>
- Ashurst. (2023, September 26). New generative AI measures in China. <https://www.ashurst.com/en/insights/new-generative-ai-measures-in-china/>
- Beijing Academy of Artificial Intelligence. (2019). Beijing AI principles. <https://www.baai.ac.cn/>
- China Briefing. (2023, July 27). How to interpret China's first effort to regulate generative AI measures. <https://www.china-briefing.com/news/how-to-interpret-chinas-first-effort-to-regulate-generative-ai-measures/>
- College of Western Idaho. (n.d.). Confucianism and Daoism. In *AI ethics* [Open textbook]. Pressbooks. <https://cwi.pressbooks.pub/aiethics/chapter/confucianism-and-daoism/>
- Coutinho, S. (2014). *An introduction to Daoist philosophies*. Columbia University Press.
- Cyberspace Administration of China, National Development and Reform Commission, Ministry of Education, Ministry of Science and Technology, Ministry of Industry and Information Technology, Ministry of Public Security, & National Radio and Television Administration. (2023). *Interim measures for the management of generative artificial intelligence services*.
- Diplo. (2025, September 23). Harmony and tools: Chinese philosophical traditions and their vision of technological change. <https://www.diplomacy.edu/blog/harmony-and-tools-chinese-philosophical-traditions-and-their-vision-of-technological-change/>
- European Parliament and 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*, L 2024/1689.
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689-707. <https://doi.org/10.1007/s11023-018-9482-5>
- Floridi, L., & Cows, 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>
- Graham, A. C. (Trans.). (2001). *Chuang-tzu: The inner chapters*. Hackett Publishing.
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Krajewska, K. (n.d.). The notion of the machine heart (ji xin 機心) in the Zhuangzi: The ethical relationship between technology and nature [Doctoral dissertation, University of Wales Trinity Saint David]. UWTSd Research Repository. <https://repository.uwtsd.ac.uk/465/>
- Lao Tzu. (1963). *Tao te ching* (D. C. Lau, Trans.). Penguin Classics.
- Musacchio, F. (2025, January 2). Wu wei: The philosophical foundation of Daoist ethics and action. https://www.fabriziomusacchio.com/weekend_stories/told/2025/2025-01-02-wu_wei/
- ÓhÉigeartaigh, S. S., Whittlestone, J., Liu, Y., Zeng, Y., & Liu, Z. (2020). Overcoming barriers to cross-cultural cooperation in AI ethics and governance. *Philosophy & Technology*, 33(4), 571-593. <https://doi.org/10.1007/s13347-020-00402-x>
- State Council of the People's Republic of China. (2017). *New generation artificial intelligence development plan*.

Balancing Innovation and Harmony: A Daoist Approach to Artificial Intelligence Law

UNESCO. (2021). Recommendation on the ethics of artificial intelligence. UNESCO Publishing.

Watson, B. (Trans.). (2013). The complete works of Zhuangzi. Columbia University Press.

Wheeler, B. (2022). The motherboard of myriad things: Daoism, Zhuangzi, and the internet. *Asian Studies*, 10(1), 425-446. <https://doi.org/10.4312/as.2022.10.1.425-446>

White & Case. (2024, July 16). Long awaited EU AI Act becomes law after publication in the EU's Official Journal. <https://www.whitecase.com/insight-alert/long-awaited-eu-ai-act-becomes-law-after-publication-eus-official-journal>

Wu, S. S. (2023). China's approach to AI regulation: Understanding Beijing's larger strategy and its global effects. Brookings Institution.

Wu, W., & Liu, S. (2023). A comprehensive review and systematic analysis of artificial intelligence regulation policies (arXiv:2307.12218). arXiv. <https://doi.org/10.48550/arXiv.2307.12218>

Yeung, K., Howes, A., & Pogrebn, G. (2020). AI governance by human rights-centred design, deliberation and oversight: An end to ethics washing. In M. D. Dubber, F. Pasquale, & S. Das (Eds.), *The Oxford handbook of ethics of AI* (pp. 77-106). Oxford University Press.

Zhu, Y. (n.d.). Integrating Daoism's Tao and Buddhism's compassion in Solo-Founder AI-Driven Nonprofit Mode (SFADNM). Sofia University.