
(When) Should We Delegate AI Governance to AIs?

Some Lessons from Administrative Law

Nicholas Caputo

Oxford Martin AI Governance Initiative
University of Oxford
Oxford OX1 3BD, United Kingdom
`nick.caputo@oxfordmartin.ox.ac.uk`

Abstract

Advanced AI systems are now being used in AI governance. Practitioners will likely delegate an increasing number of tasks to them as they improve and governance becomes harder. However, using AI for governance risks serious harms because human practitioners may not be able to understand AI decisions or determine whether they are aligned to the user’s interests. Delegation may also undermine governance’s legitimacy. This paper begins to develop a principled framework for when to delegate AI governance to AIs and when (and how) to maintain human participation. Administrative law, which governs agencies that are (1) more expert in their domains than the legislatures that create them and the courts that oversee them and (2) potentially misaligned to their original goals, offers useful lessons. Administrative law doctrine provides examples of clear, articulated rules for when delegation can occur, what delegation can consist of, and what processes can keep agencies aligned even as they are empowered to achieve their goals. The lessons of administrative law provide a foundation for how AI governance can use AI in a safe, accountable, and effective way.

1 Introduction

Frontier AI systems continue to improve in both their capabilities and the scope of tasks they can accomplish [38, 43]. Because these systems are dual use, these improvements come with risks, especially as AIs become increasingly agentic and self-directed [18]. To keep up with these novel risks, regulation of advanced AI will have to incorporate AI tools. New technologies that create problems have always been met with technical innovations in governance, from the railroad and telegraph enabling both new kinds of torts [16] and new ways of coordinating officials [48] to social media’s harms and the development of algorithmic content moderation [35, 29].

However, advanced AI presents a special puzzle for governance. Unlike the telegraph or classifier algorithms, these new systems can participate in the design of their own governance institutions because they are capable of the kind of research, analysis, and writing that constitute governance itself. Indeed, everyone working in AI governance, including within frontier companies [51], likely already relies on chatbots at least for tasks like research assistance and editing. If frontier AIs continue to improve, AI-designed governance institutions will likely someday outperform those designed by humans. But, especially given concerns around democratic representation in governance and AI misalignment, simply delegating AI governance decisions to AI without considering the risks of doing so and establishing clear rules for delegation would be dangerous.

We must develop a principled framework for when decisions about how to govern AI should be delegated to AIs. Fortunately, governance has faced a similar problem before. Administrative

agencies are expert bodies created by legislatures that delegate power to them to solve specialized problems. They are often empowered to make rules, adjudicate cases, investigate potential violations, and generally act with substantial independence [50, 2]. Within their domains, they are more expert than the legislatures that created them and the courts provide judicial review of their actions, creating an “information asymmetry” [31, 28]. They also sometimes experience “preference divergence” or misalignment from their public purpose, either through capture or internal pressures [40]. Legislatures and courts must therefore find ways to simultaneously empower and oversee agencies while facing a deficit of information and competence compared to the agencies themselves. They seek to resolve this problem through the tools of administrative law.

This paper begins to lay out a framework for delegating governance decisions to AI, drawing on historical lessons from administrative law. It uses five key parts of administrative law doctrine, from nondelegation to “hard look” review, to illustrate how Congress and the courts navigate the problem of delegating authority to expert agencies. In short, it suggests that careful, limited delegation bounded by review processes can help navigate the tradeoff between empowering agencies and risking misalignment. Our analysis is preliminary, sketching commonalities without providing complete guides for delegation, but it suggests that there are useful lessons to draw for this critical governance question.

2 Background

Substantial literature exists on whether and when to delegate decisions to AI systems. Researchers have investigated bail and bond determinations in the criminal justice system [14, 34], child protection hearings [27], tenancy determinations [33], credit provision [42, 32, 17], and similar areas. This research has demonstrated the many problems that such delegation can cause, including replicating bias [19], invading privacy [56], entrenching inequalities [27], and causing various dignitary harms [21]. However, the existing literature also suggests that in some cases algorithms are better decisionmakers than humans. For example, algorithms may improve on biased or inattentive human judges in bail determinations [34], and recent research has shown that AI systems can make better diagnoses than many doctors [26]. Furthermore, algorithms are often (though not always) more explainable and corrigible than humans, so once a problem has been identified in an algorithm, it can be rectified in a way that might not be possible for humans [30].

An even vaster administrative law literature covers delegation, deference, and how to create expert governance while preserving accountability and legitimacy. Administrative law begins from the premise that Congress can authorize the creation of administrative agencies empowered to make rules and adjudicate cases in pursuit of specified objectives [52, 1]. However, administrative power is sharply limited by law, which ensures that agencies act in democratic and responsible ways even where they are more expert and capable than Congress and courts [50]. Some scholars have begun to focus on the use of AI in administration and how the law can help ensure its legitimacy. One key question is whether the use of AI by agencies can meet the requirements of key administrative doctrines including nondelegation, procedural due process, and equal protection [21, 22, 49]. Broader analyses consider whether machine learning applications undermine fundamental attributes like the accountability and legitimacy of administration [24, 20] and how courts might respond [23]. However, because AI systems that might participate in governance design itself are new, the scholarship has neglected how administrative law could inform AI delegation, focusing instead on how AI might change administrative law.

3 Administrative Law’s Lessons

Some key lessons emerge from the administrative law literature that might be useful for AI delegation. This list is not exhaustive, but rather illustrative of how inspiration could be drawn.

First, administrative law shows it’s possible to navigate tradeoffs between capabilities and alignment similar to those facing AI governance. Administrative agencies must be empowered to perform their governance functions. If Congress or the courts intervene too much in administrative decision-making, especially in areas that are highly technical and specialized, the benefits of delegation will be lost. At the same time, simply letting agencies define and pursue their own objectives without oversight risks them acting contrary to their intended purpose [53]. The recent debate about *Chevron* deference

[8] concerned this question, and the Supreme Court decided that deferring to agency interpretations of their own authorizing statutes gave agencies too much autonomy with too little oversight [11]. As a field, administrative law aims to fine-tune the balance between capabilities and alignment of agencies as they act. AI governance already faces similar problems of "information asymmetry" and "preference divergence" to administrative law. It currently lacks, however, the tools to build oversight that would allow us to pick the optimal spot on the tradeoff curve, as administrative law aims to do.

Second, delegation should be done for specific and limited purposes, and certain areas might be best marked out as non-delegable. Congress' power to delegate authority to administrative agencies is limited. It must provide an "intelligible principle" that guides the agency's action rather than simply providing a grant of authority [1]. Furthermore, for so-called "major questions" of significant importance, Congress must be extremely explicit about its delegation [9, 10]. Under the Constitution, delegation should be a limited tool that supplements Congress' ability to perform its functions rather than a full grant of power to an external body [52]. Governance delegations to AIs might need to be similarly limited in scope and context. Existing AI laws like the EU AI Act already distinguish the kinds of functions AI systems can perform [45]. Extending these concepts to governance, we might decide that key mechanisms like critical risk thresholds or the use of autonomous AI in the military must be designed by humans or by AIs with significant human instruction and oversight because of their significance for human life.

Third, certain process requirements can help compensate for a lack of substantive understanding of why an agent is deciding in a particular way. Because administrative agencies are expert bodies operating in specialized fields, other branches of government struggle to provide oversight [37]. Indeed, Congress' inability to understand what is happening in a given domain is often a key reason it creates an administrative agency in the first place [25]. Courts providing judicial review similarly find it difficult to provide substantive review of agency decisions that they may lack the context to understand [8]. But this lack of understanding does not mean that oversight is impossible. Administrative law uses procedural requirements placed on agencies to make up for the lack of substantive knowledge. For example, courts can check whether agencies met the requirements that they give reasons for their actions and consider the whole record of evidence when deciding, even if they cannot judge the correctness of the ultimate decisions [7]. Tools like "hard look review" provide sufficient oversight of administrative agencies to keep them accountable and ensure the legitimacy of their actions. Agencies are required to give the true reasons why they took their actions, not just make up plausible justifications after the fact [4]. Requiring that AI systems similarly provide reasons for their decisions and follow established decision-making processes would help people understand the governance decisions they make or at least provide a partial guarantee that the choices were made in a pro-social way. Interpretability, chain-of-thought monitoring, and similar techniques [47, 54, 36] could even provide a more complete guarantee of legitimacy in the AI context than we get in the context of human administrative agencies.

Fourth, public participation in decision-making enhances both its quality and its legitimacy. Administrative agencies undergo significant public process when making new rules and adjudicating cases. The "notice and comment" process ensures that concerned parties are notified and given the chance to participate in shaping or stopping actions that might affect them [3, 6]. Adjudications must meet various procedural due process requirements that seek to guarantee that no rights are violated in deciding the case [5]. Furthermore, public participation provides useful new inputs to decision-making that improve how agencies carry out their work. Those affected by an agency action likely have special local knowledge that would be useful for the agency to consider when making its decision [50]. Including public voices provides the agencies with new information that might change what they think is the best course of action. Delegation of AI governance should also include public participation in the rule-making process. AIs making decisions about governance should be required to take human input into account and address it in comprehensive, rational, and rule-bound ways. Such participation would ensure some degree of legitimacy in the process [55] and also ensure that human voices and perspectives shaped how AI is governed.

Finally, administrative law teaches that it is not always a bad idea to rely on the comparative expertise and competence of others. Congress would have to entirely reconstruct itself to serve the variety of functions that the administrative state currently performs, and in doing so, it would undermine its own purpose as a representative legislature. The complex problems of modern life require sophisticated forms of governance that come with tradeoffs, but administrative law and similar institutions can help us make better decisions about how to handle those tradeoffs [37]. AI will have to be a part of the

solution to AI's problems. The challenges of AI will require new forms of government just as the new problems created by technological progress in the early twentieth century demanded the creation of the administrative state [39]. In fact, AI might enable better governance than is possible with human-run agencies. The capabilities/alignment tradeoff that governance of administrative agencies faces has been relatively immutable over time, a product of human psychology and small advances in technology. In AI, however, advances in the fields of capabilities and alignment, giving people greater insight into why AI systems act than is possible with humans, may create the opportunity for more effective and responsible government. Combining the technical promise of AI with the insights of administrative law could open up powerful new opportunities for governance.

4 Example Applications

First, AI risk managers must set thresholds for when to conduct comprehensive evaluations of the risks presented by their systems [45, 44, 15]. Setting thresholds requires understanding system capabilities, the external risk environment, and how different risks can emerge across a variety of domains from biology to cybersecurity [13]. Administrative law suggests that delegating threshold setting could be a reasonable response to this complexity as long as sufficiently clear ("intelligible principle"-type) guidance is given to the AI and process requirements are followed. The ultimate decision to accept or reject the proposed threshold might still need to rest with an accountable human, but the technical details of research and implementation could be offloaded to some extent.

Second, AI companies must allocate limited red-teaming resources to evaluate the safety of their systems [12]. Determining how to deploy those resources is an optimization problem that requires synthesizing huge quantities of internal and external data on where risks might originate. This scenario might be one in which delegation to an AI is reasonable. AIs might be more competent at the relevant kinds of analysis and optimization than humans are, as well as better able to overcome their biases about where risks might be coming from. Administrative law-style checks like requiring explanations for resource allocation and human sign-off on huge shifts would help ensure that red-teaming stays on track and aligned with the overall goals of the company.

Third, AI companies are already creating mechanisms that determine how much compute to allocate different users to solve their problems. Some allocation is simply done by pricing tiers, but consumer apps like ChatGPT now incorporate "routers" for allocation [41]. Access to compute might someday become an essential part of modern life, so control over who can access it and under what circumstances will be highly significant. Administrative agencies must follow procedural rules when distributing essential benefits because of how important they are for people's lives [46]. It may soon become necessary for AI companies or governments to similarly implement rules governing how compute is distributed to people and laying out whether and how users might have rights to compute that they can enforce against those seeking to limit or rescind access. Administrative law's doctrines could provide inspiration for ways to protect people in the novel contexts of AI, drawing on old lessons for these new problems.

5 Limitations

The foregoing analysis is preliminary and not meant to provide final rules for applying administrative law's lessons to AI governance. Both fields are deep and rapidly changing, as AI develops and the Supreme Court reconsiders foundational parts of administrative law doctrine. There are many areas where administrative law does not fit the kinds of problems that AI governance is facing. This paper is intended to illustrate the kinds of connections that implicitly exist between law and AI and start a conversation drawing on them to help improve each. Further work formalizing these insights through the creation of decision rules for delegation and the like would represent a helpful path forward.

6 Conclusion

Administrative law has its challenges and is currently undergoing a period of serious upheaval. Even at its best, it is an imperfect set of often vague rules trying to handle the massive challenges of administrative governance. Yet this imperfect body of doctrine has guided decades of controversial

delegation. AI governance can learn key lessons of governance design from administrative law, and determining when and how to use AI in governance is one important place to start.

References

- [1] J.W. Hampton, Jr. & Co. v. United States. 276 U.S. 394, 1928.
- [2] Humphrey’s Executor v. United States. 295 U.S. 602, 1935.
- [3] Administrative Procedure Act. P.L.79-404 60 Stat. 237, 1946.
- [4] SEC v. Chenery Corp. (Chenery II). 332 U.S. 194, 1947.
- [5] Mathews v. Eldridge. 424 U.S. 319, 1976.
- [6] United States v. Nova Scotia Food Products Corp. 568 F.2d 240 (2d Cir. 1977), 1977.
- [7] Motor Vehicle Mfrs. Ass’n v. State Farm Mutual Automobile Ins. Co. 463 U.S. 29, 1983.
- [8] Chevron U.S.A., Inc. v. Natural Resources Defense Council, Inc. 467 U.S. 837, 1984.
- [9] FDA v. Brown & Williamson Tobacco Corp. 529 U.S. 120, 2000.
- [10] West Virginia v. EPA. 597 U.S. 697, 2022.
- [11] Loper Bright Enterprises v. Raimondo. 603 U.S. 369, 2024.
- [12] Lama Ahmad, Sandhini Agarwal, Michael Lampe, and Pamela Mishkin. Openai’s approach to external red teaming for ai models and systems, 2025.
- [13] Markus Anderljung, Joslyn Barnhart, Anton Korinek, Jade Leung, Cullen O’Keefe, Jess Whittlestone, Shahar Avin, Miles Brundage, Justin Bullock, Duncan Cass-Beggs, Ben Chang, Tatum Collins, Tim Fist, Gillian Hadfield, Alan Hayes, Lewis Ho, Sara Hooker, Eric Horvitz, Noam Kolt, Jonas Schuett, Yonadav Shavit, Divya Siddarth, Robert Trager, and Kevin Wolf. Frontier ai regulation: Managing emerging risks to public safety, 2023.
- [14] Julia Angwin, Jeff Larson, Surya Mattu, and Lauren Kirchner. Machine bias: There’s software used across the country to predict future criminals. and it’s biased against blacks. *ProPublica*, May 2016.
- [15] Anthropic. Responsible scaling policy (version 2.2), May 2025.
- [16] Evelyn Atkinson. Telegraph torts: The lost lineage of the public service corporation. *Michigan Law Review*, 121(8):1365, 2023.
- [17] Robert Bartlett, Adair Morse, Richard Stanton, and Nancy Wallace. Consumer-lending discrimination in the fintech era. *Journal of Financial Economics*, 143(1):30–56, 2022.
- [18] Yoshua Bengio, Geoffrey Hinton, Andrew Yao, Dawn Song, Pieter Abbeel, Trevor Darrell, Yuval Noah Harari, Ya-Qin Zhang, Lan Xue, Shai Shalev-Shwartz, Gillian Hadfield, Jeff Clune, Tegan Maharaj, Frank Hutter, Atılım Güneş Baydin, Sheila McIlraith, Qiqi Gao, Ashwin Acharya, David Krueger, Anca Dragan, Philip Torr, Stuart Russell, Daniel Kahneman, Jan Brauner, and Sören Mindermann. Managing extreme ai risks amid rapid progress. *Science*, 384(6698):842–845, May 2024.
- [19] Joy Buolamwini and Timnit Gebru. Gender shades: Intersectional accuracy disparities in commercial gender classification. In *Proceedings of the First Conference on Fairness, Accountability and Transparency*, pages 77–91, 2018.
- [20] Ryan Calo and Danielle Keats Citron. The automated administrative state: A crisis of legitimacy. *Emory Law Journal*, 70(4):797, 2021.
- [21] Danielle Keats Citron. Technological due process. *Washington University Law Review*, 85(5):1249, 2008.
- [22] Cary Coglianese and David Lehr. Regulating by robot: Administrative decision making in the machine-learning era. *Georgetown Law Journal*, 105:1147, 2017.

- [23] Ashley Deeks. The judicial demand for explainable artificial intelligence. *Columbia Law Review*, 119(7):1829, 2019.
- [24] David Freeman Engstrom and Daniel E. Ho. Algorithmic accountability in the administrative state. *Yale Journal on Regulation*, 37:800, 2020.
- [25] David Epstein and Sharyn O’Halloran. *Delegating Powers: A Transaction Cost Politics Approach to Policy Making Under Separate Powers*. Cambridge University Press, 1999.
- [26] Andre Esteva, Brett Kuprel, Roberto A. Novoa, Justin Ko, Susan M. Swetter, Helen M. Blau, and Sebastian Thrun. Dermatologist-level classification of skin cancer with deep neural networks. *Nature*, 542(7639):115–118, 2017.
- [27] Virginia Eubanks. *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor*. St. Martin’s Press, 2018.
- [28] Jacob Gersen. Designing agencies. In Daniel A. Farber and Anne Joseph O’Connell, editors, *Research Handbook on Public Choice and Public Law*, pages 333–362. Edward Elgar Publishing, Cheltenham, 2010.
- [29] Tarleton Gillespie. *Custodians of the Internet: Platforms, Content Moderation, and the Hidden Decisions That Shape Social Media*. Yale University Press, 2018.
- [30] Sharad Goel, Ravi Shroff, Jennifer Skeem, and Christopher Slobogin. The accuracy, equity, and jurisprudence of criminal risk assessment. In Roland Vogl, editor, *Research handbook on big data law*. Edward Elgar Publishing, 2021.
- [31] John D. Huber and Charles R. Shipan. Politics, delegation, and bureaucracy. In Barry R. Weingast and Donald A. Wittman, editors, *The Oxford Handbook of Political Economy*, pages 256–272. Oxford University Press, Oxford, 2006.
- [32] Mikella Hurley and Julius Adebayo. Credit scoring in the era of big data. *Yale Journal of Law and Technology*, 18:148, 2016.
- [33] Lauren Karpinski. The discriminatory impacts of ai-powered tenant screening programs, 2025.
- [34] Jon Kleinberg, Himabindu Lakkaraju, Jure Leskovec, Jens Ludwig, and Sendhil Mullainathan. Human decisions and machine predictions. *The Quarterly Journal of Economics*, 133(1):237–293, 2018.
- [35] Kate Klonick. The new governors: The people, rules, and processes governing online speech. *Harvard Law Review*, 131:1598, 2017.
- [36] Tomek Korbak, Mikita Balesni, Elizabeth Barnes, Yoshua Bengio, Joe Benton, Joseph Bloom, Mark Chen, Alan Cooney, Allan Dafoe, Anca Dragan, Scott Emmons, Owain Evans, David Farhi, Ryan Greenblatt, Dan Hendrycks, Marius Hobbhahn, Evan Hubinger, Geoffrey Irving, Erik Jenner, Daniel Kokotajlo, Victoria Krakovna, Shane Legg, David Lindner, David Luan, Aleksander Mądry, Julian Michael, Neel Nanda, Dave Orr, Jakub Pachocki, Ethan Perez, Mary Phuong, Fabien Roger, Joshua Saxe, Buck Shlegeris, Martín Soto, Eric Steinberger, Jasmine Wang, Wojciech Zaremba, Bowen Baker, Rohin Shah, and Vlad Mikulik. Chain of thought monitorability: A new and fragile opportunity for ai safety, 2025.
- [37] James M. Landis. *The Administrative Process*. Yale University Press, 1938.
- [38] Nestor Maslej, Loredana Fattorini, Raymond Perrault, Vanessa Parli, Anka Reuel, Erik Brynjolfsson, John Etchemendy, Katrina Ligett, Terah Lyons, James Manyika, Juan Carlos Niebles, Yoav Shoham, Russell Wald, and Jack Clark. Artificial intelligence index report 2024, 2024.
- [39] Thomas K. McCraw. *Prophets of Regulation: Charles Francis Adams, Louis D. Brandeis, James M. Landis, Alfred E. Kahn*. Harvard University Press, Cambridge, MA, 1984.
- [40] William A. Niskanen. *Bureaucracy and Representative Government*. Aldine-Atherton, 1971.
- [41] Doug OLaughlin, Dylan Patel, Wei Zhou, and AJ Kourabi. GPT-5 Set the Stage for Ad Monetization and the SuperApp. Blog post, August 2025.

- [42] Cathy O’Neil. *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown Publishing Group, 2016.
- [43] OpenAI. Introducing gpt-5. Blog post, August 2025.
- [44] OpenAI. Preparedness framework (v2), April 2025.
- [45] Regulation (EU) 2024/1689 of the European Parliament and of the Council. 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). Official Journal of the European Union, OJ L, 2024/1689, 12.7.2024, jul 2024.
- [46] Charles A. Reich. The New Property. *Yale Law Journal*, 73(5):733, 1964.
- [47] Cynthia Rudin. Stop explaining black box models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5):206–215, 2019.
- [48] James C. Scott. *Seeing Like a State: How Certain Schemes to Improve the Human Condition Have Failed*. Yale University Press, 1998.
- [49] Alicia Solow-Niederman. Administering artificial intelligence. *Southern California Law Review*, 93:633, 2020.
- [50] Richard B. Stewart. The reformation of american administrative law. *Harvard Law Review*, 88(8):1667, 1975.
- [51] Charlotte Stix, Matteo Pistillo, Girish Sastry, Marius Hobbhahn, Alejandro Ortega, Mikita Balesni, Annika Hallensleben, Nix Goldowsky-Dill, and Lee Sharkey. Ai behind closed doors: a primer on the governance of internal deployment, 2025.
- [52] Peter L. Strauss. The place of agencies in government: Separation of powers and the fourth branch. *Columbia Law Review*, 84(3):573, 1984.
- [53] Adrian Vermeule. The Deference Dilemma. *George Mason Law Review*, 31(2):619, 2024.
- [54] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models. In *NIPS’22: Proceedings of the 36th International Conference on Neural Information Processing Systems*, 2022.
- [55] Daniel Wilf-Townsend and Kevin Tobia. Ai-generated legal texts. SSRN Working Paper, May 2025.
- [56] Shoshana Zuboff. *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. Profile Books, 2019.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract and introduction clearly state the claims made, the framework of analysis, and the key contributions the paper is intended to make.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The paper specifically discusses the limitations of attempting to draw lessons from administrative law for AI, in the "Limitations" section and throughout.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: The paper does not include theoretical results.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[NA\]](#)

Justification: The paper does not include experiments.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[NA\]](#)

Justification: The paper does not include experiments requiring code.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[NA\]](#)

Justification: The paper does not include experiments.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[NA\]](#)

Justification: The paper does not include experiments.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[NA\]](#)

Justification: The paper does not include experiments.

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: As a theoretical contribution from law, the paper does not implicate any ethical concerns.

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: The paper considers how the delegation of governance to AI systems might cause severe harms and undermine the extent of human participation in governance.

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper poses no such risks.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: The paper does not use existing assets.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components. We have only used LLMs for editing and formatting.