



Algorithmic Accountability in Technology-Enabled Welfare Administration

Dr. Rachel Thompson

Professor, University of Oxford, United Kingdom

Abstract

Technology-enabled welfare administration increasingly relies on data integration, automated eligibility screening, risk classification, fraud detection, predictive analytics, and artificial intelligence-assisted decision support. These systems can improve administrative speed, consistency, targeting, and service coordination, but their use in decisions affecting income support, housing assistance, disability benefits, child welfare, unemployment protection, and other essential services creates a distinctive accountability problem. Administrative responsibility may become distributed across software developers, private vendors, public agencies, data infrastructures, frontline officers, and automated decision processes, making it difficult for affected individuals to determine why a decision occurred or who is responsible for correcting it.

Contemporary public-sector governance therefore requires algorithmic accountability to extend beyond technical transparency toward institutional responsibility, procedural fairness, traceability, meaningful review, and accessible redress. OECD evidence indicates that artificial intelligence is already used in at least one governmental area in 35 of 36 surveyed OECD countries, while higher-stakes uses continue to demand stronger transparency, data quality, and assurance mechanisms.

This paper develops an integrated accountability framework for technology-enabled welfare administration through a structured conceptual synthesis of public-administration scholarship, automated decision-making research, AI governance standards, and regulatory developments. The framework organizes accountability around seven interdependent dimensions: traceability, explainability, bias testing, meaningful human review, appeal and redress, data governance, and independent auditability. Rather than treating a human official's presence as sufficient protection, the study argues that accountability must operate throughout the administrative and technological lifecycle.

An illustrative maturity analysis demonstrates how a welfare agency could evaluate these dimensions before and after implementing the proposed framework. The paper concludes that legitimate welfare automation depends not merely on whether an algorithm is accurate, but on whether public institutions can justify its use, monitor its consequences, explain consequential decisions, identify responsible actors, and provide effective mechanisms for correction.

Keywords: algorithmic accountability; welfare administration; automated decision-making; artificial intelligence; public administration; administrative justice; algorithmic transparency; social protection



1. Introduction

Digital transformation has altered the administrative infrastructure through which governments identify beneficiaries, verify eligibility, distribute assistance, investigate irregularities, and manage public resources. Welfare agencies increasingly work with interconnected databases, automated workflows, predictive models, digital identity systems, risk-scoring techniques, and AI-supported decision tools. Such technologies can reduce repetitive administrative work and potentially make public services more responsive, particularly where agencies manage large numbers of applications under limited staffing and budget conditions. OECD reporting in 2026 shows that government use of AI has become widespread across its surveyed member countries, although adoption remains uneven and the most consequential policy and oversight applications receive greater scrutiny because of their implications for transparency, assurance, and data quality.

The welfare context makes this transformation especially significant because administrative decisions can immediately affect material security. An erroneous eligibility classification may delay income assistance; a risk score may initiate an investigation; a data-matching error may interrupt a benefit; and an automated recommendation may influence how a caseworker evaluates an applicant. These consequences distinguish welfare algorithms from many ordinary commercial recommendation systems. Levy, Chasalow, and Riley observe that algorithmic decision systems are used across public-sector domains including benefits provision and emphasize their implications for accountability, privacy, inequality, and participation. The United Nations' work on digital technology and social protection has likewise drawn attention to the human-rights implications created when welfare systems increasingly depend on automated data processing.

The central governance difficulty is that algorithmic administration can disperse responsibility. A frontline officer may rely on an output produced by a system that was procured by another department, designed by an external vendor, trained or calibrated using historical administrative data, and integrated with databases controlled by separate public institutions. When an individual experiences an adverse outcome, identifying the origin of the decision may therefore become difficult. Technical opacity is only one dimension of the problem. Equally important are organizational opacity, unclear procurement responsibilities, limited documentation, inadequate staff authority to challenge automated recommendations, weak audit arrangements, and inaccessible appeal procedures. OECD's AI principles accordingly connect accountability with traceability across datasets, processes, and decisions throughout the AI lifecycle rather than treating accountability as a single disclosure event.

Regulatory developments increasingly recognize the high stakes of algorithmic public administration. The European Union's AI Act establishes obligations for high-risk AI systems involving accuracy, robustness, documentation, oversight, registration, monitoring, and information to persons affected by certain AI-assisted decisions. The consolidated legislation also explicitly warns against automation bias and requires deployers to remain capable of understanding and appropriately interpreting system outputs.

These requirements illustrate a broader movement from abstract ethical principles toward operational accountability obligations. NIST's AI Risk Management Framework similarly treats governance and trustworthiness as matters that should be incorporated into the design, development, deployment, use, and evaluation of AI systems rather than addressed only after a harmful event occurs.



Against this background, this study conceptualizes algorithmic accountability as an institutional capacity to justify, document, examine, contest, correct, and learn from technology-mediated welfare decisions. The paper does not claim results from a completed beneficiary survey or administrative experiment. Instead, it develops a conceptual-methodological framework grounded in scholarly literature and recognized governance standards. Its contribution lies in connecting technical safeguards with administrative-law values and operational welfare procedures, thereby shifting the central question from whether a human remains somewhere in the decision process to whether the entire sociotechnical decision system remains answerable to affected citizens and competent oversight institutions.

2. Objectives of the Study

2.1 Developing an Integrated Accountability Architecture

The first objective is to identify the institutional and technical elements required for accountable algorithm-assisted welfare administration. Accountability is treated as a multidimensional governance architecture rather than a synonym for transparency. The study therefore examines traceability, explanation, bias monitoring, human review, data stewardship, contestability, and audit arrangements as mutually reinforcing safeguards.

2.2 Connecting Automation With Procedural Justice

The second objective is to examine how technological efficiency can be reconciled with the principles traditionally associated with legitimate public administration. These include consistency, equal treatment, reason-giving, reviewability, proportionality, and the opportunity to challenge a materially adverse administrative decision. Research on automated administration has increasingly connected responsible system design with established expectations of good administration, including transparency, equality, and accountability.

2.3 Establishing an Operational Governance Model

The third objective is to translate broad responsible-AI principles into a usable governance model for welfare agencies. The proposed approach identifies controls that can be incorporated before procurement, during model development and deployment, at the point of administrative decision-making, and during post-deployment monitoring. This lifecycle orientation reflects both OECD traceability principles and NIST's emphasis on continuous AI risk management.

3. Review of Literature

3.1 Algorithms, Welfare Administration, and Public-Sector Responsibility

Levy, Chasalow, and Riley's analysis of algorithmic decision-making in government demonstrates that algorithmic systems cannot be evaluated only through predictive performance. Their effects depend on problem definition, acquisition, implementation, administrative context, and evaluation. They further emphasize that data limitations and institutional choices influence the fairness and legitimacy of public-sector algorithms. This approach is particularly relevant to welfare administration, where historical administrative records may reflect differences in reporting, enforcement intensity, access to services, bureaucratic discretion, and institutional priorities.



Digital welfare systems also raise questions about how governments conceptualize claimants. Technology can simplify access and accelerate payments, but intensive data matching and risk classification can simultaneously transform beneficiaries into objects of continuous administrative assessment.

The World Bank recognizes that digital technologies can enhance social-protection access and monitoring while also generating privacy, exclusion, and bias risks when governance frameworks and institutional capacity are insufficient. Accordingly, the policy objective should not be either unconditional automation or complete rejection of technology. The more productive question concerns the conditions under which technological systems strengthen rather than weaken legitimate welfare administration.

Recent scholarship reinforces the need for context-sensitive safeguards. Agbabiaka and colleagues' systematic review of AI-enabled automated decision-making in the public sector identifies a broad set of trustworthiness requirements, demonstrating that responsible automation involves more than model accuracy. The literature therefore increasingly treats government algorithms as sociotechnical systems in which institutional structures, people, data, policies, software, and organizational incentives jointly influence outcomes.

3.2 Limits of Transparency and Human Oversight

A common accountability response is to require a human decision-maker to remain involved. Although this requirement is intuitively attractive, research shows that human presence alone cannot guarantee meaningful control. Green's analysis of government policies requiring human oversight argues that institutional actors may rely excessively on frontline human reviewers without establishing whether those reviewers possess sufficient information, authority, time, independence, or cognitive capacity to challenge algorithmic recommendations. Green therefore proposes greater emphasis on institutional rather than merely individual oversight.

This concern is directly relevant to welfare casework. Sztandar-Sztanderska's research on automated decision-making in welfare administration examines how contextual conditions shape the ability of frontline employees to supervise automated systems. The work demonstrates why oversight cannot be assessed independently of organizational circumstances surrounding caseworkers' actual responsibilities and working environments. A nominal right to override an automated recommendation has limited value if staff cannot understand the model, are under strong productivity pressure, lack access to relevant evidence, or believe departures from the system will require burdensome justification. Meaningful oversight must therefore contain at least three elements.

First, reviewers need sufficient information to understand the role and limitations of the system. Second, they need genuine authority to reconsider or override its recommendation. Third, institutional arrangements must prevent the automated output from becoming the unchallenged default. The EU AI Act reflects this concern by requiring attention to automation bias in high-risk systems and by emphasizing users' capacity to interpret outputs correctly.

3.3 Accountability as Lifecycle Governance

The strongest emerging governance approaches treat accountability as continuous. OECD principles require AI actors to maintain traceability regarding datasets, processes, and decisions and to apply systematic risk management throughout the AI lifecycle. NIST similarly frames responsible AI governance around activities extending from system design and development to use and evaluation.

The EU AI Act adds a regulatory dimension to this lifecycle model through requirements concerning risk management, documentation, human oversight, accuracy, registration, and post-market monitoring for covered high-risk systems. The legislation requires certain public-authority deployments to be registered and provides for continuous collection and analysis of performance information after deployment. This represents an important conceptual shift: a system cannot be regarded as accountable merely because it passed an initial technical test.

The literature nevertheless remains fragmented. Technical research often emphasizes explainability, fairness metrics, or model validation, while public-administration scholarship emphasizes discretion, legitimacy, and institutional responsibility. Legal approaches focus on rights and procedural safeguards, whereas operational guidance focuses on risk-management processes. The research gap is therefore not the absence of accountability principles, but insufficient integration of these principles into a welfare-specific administrative architecture connecting technology, organizational responsibility, claimant rights, and continuous oversight.

4. Research Methodology

4.1 Research Design and Evidence Base

The study adopts a structured conceptual-synthesis design. It is not presented as a completed population survey, randomized experiment, or administrative dataset analysis. The analytical corpus combines peer-reviewed literature on public-sector automated decision-making with authoritative institutional materials addressing AI governance, public administration, social protection, and high-risk automated systems.

Evidence was identified through targeted searches conducted up to August 12, 2026 using combinations of the terms algorithmic accountability, automated decision-making, welfare administration, public-sector AI, human oversight, algorithmic transparency, social protection technology, AI risk management, and high-risk public services.

Priority was given to peer-reviewed scholarship and primary institutional sources including the OECD, NIST, European Union, United Nations, and World Bank. Materials were included when they directly addressed public-sector algorithms, welfare or social-protection technology, accountability safeguards, automated administrative decisions, or recognized governance requirements. Purely commercial AI applications with no meaningful public-administration relevance were excluded.

4.2 Analytical Coding Framework

The evidence was synthesized using seven accountability dimensions derived from recurring governance concerns across the literature and official frameworks. The analysis focused not simply on whether each dimension existed, but on its administrative purpose and the institutional actor responsible for operationalizing it.

Table 1: Analytical Dimensions of Algorithmic Accountability

Accountability Dimension	Core Administrative Question	Primary Safeguard
Traceability	Can the agency reconstruct how a decision was produced?	Decision logs, data lineage, model/version records
Explainability	Can the affected person understand the material basis of the decision?	Plain-language reasons and relevant decision factors
Bias Testing	Are unequal or systematically harmful outcomes being detected?	Pre-deployment and periodic subgroup testing
Meaningful Human Review	Can qualified officials independently reconsider the automated output?	Override authority, training, time, and supporting evidence
Appeal and Redress	Can an affected person challenge and correct the decision?	Accessible review, correction, and remedy procedures
Data Governance	Are data appropriate, accurate, lawful, proportionate, and controlled?	Data-quality rules, access controls, retention policies
Auditability	Can an independent body evaluate the system and its operation?	Internal and external audits, documentation, monitoring

4.3 Framework Construction and Illustrative Validation

The proposed framework was constructed by mapping each accountability dimension to four stages of the welfare-technology lifecycle: authorization and procurement, development and testing, operational decision-making, and monitoring/redress. This approach prevents agencies from concentrating accountability solely at the final decision stage.

For analytical illustration, a five-point maturity scale was created in which 0 represents the absence of an accountability mechanism and 5 represents a strongly institutionalized mechanism with documented responsibility, regular evaluation, and enforceable review processes. The numerical profile presented later is deliberately simulated. It does not represent observed data from a real welfare agency and is included only to demonstrate how the proposed framework could support institutional self-assessment.

5. Analysis

5.1 Accountability Risks Across the Welfare Decision Lifecycle

The analysis indicates that accountability failures may originate before an algorithm processes any individual case. During procurement, agencies may acquire proprietary systems without securing adequate documentation, audit rights, performance-disaggregation requirements, or access to

information necessary for explaining decisions. A procurement contract that protects vendor secrecy more strongly than administrative review can create structural opacity even when frontline officers act responsibly.

Development and testing introduce a second risk layer. Welfare records are administrative artifacts rather than neutral representations of social reality. Historical data may contain missing information, different levels of institutional scrutiny across communities, policy changes, reporting inconsistencies, and proxy variables associated with socioeconomic disadvantage. The purpose of bias testing is therefore not to assert that every statistical difference constitutes discrimination; it is to require agencies to investigate whether systematic differences are justified, legally permissible, and consistent with the welfare program's objectives.

At the operational stage, the principal risks involve automation bias, inadequate explanation, and diffusion of responsibility. An official may technically retain authority while practically treating the algorithmic recommendation as presumptively correct. The EU framework expressly recognizes the danger of automatically or excessively relying on high-risk AI outputs. Green's research similarly demonstrates why nominal human involvement cannot substitute for institutional safeguards.

Post-deployment accountability is equally important because model performance and administrative environments can change. Updated policies, economic conditions, demographic patterns, data sources, and operational practices may alter how a system performs after initial validation. The EU AI Act's post-market monitoring approach illustrates the importance of systematically collecting and analyzing performance information throughout the operational life of high-risk systems.

5.2 Proposed Welfare Algorithm Accountability Matrix

Table 2: Lifecycle Accountability Matrix for Technology-Enabled Welfare Administration

Lifecycle Stage	Principal Risk	Required Accountability Control	Responsible Actor
Authorization	Technology adopted without demonstrated necessity	Algorithmic impact and proportionality assessment	Agency leadership / oversight authority
Procurement	Vendor opacity and unclear responsibility	Contractual documentation, testing, and audit rights	Procurement authority
Data Preparation	Inaccurate, incomplete, or historically distorted records	Data-quality and representativeness assessment	Data governance team

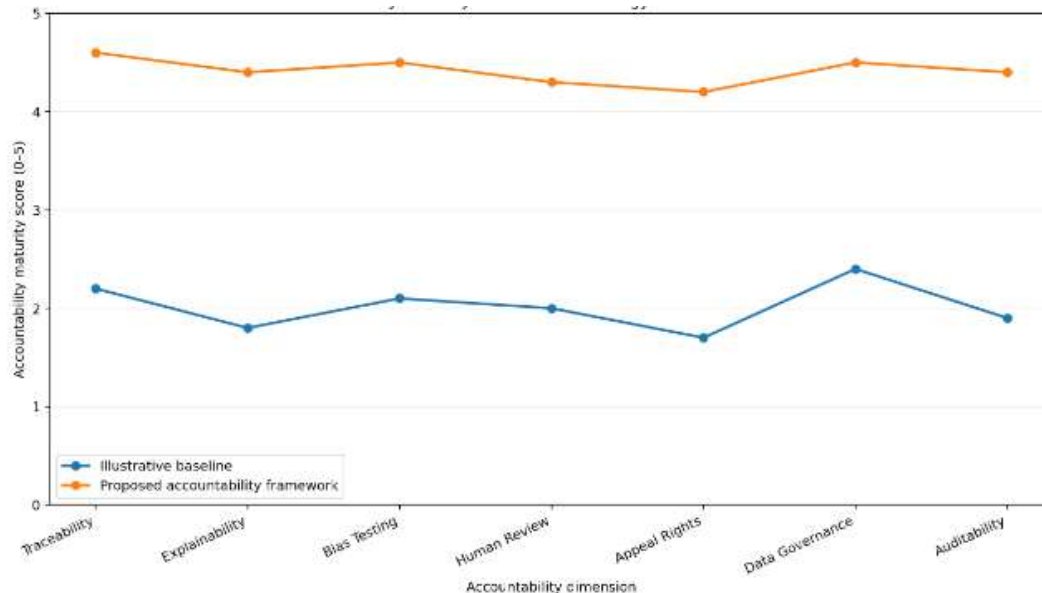
Lifecycle Stage	Principal Risk	Required Accountability Control	Responsible Actor
Model Development	Unfair performance or unsuitable objective function	Validation, subgroup testing, error analysis	Technical team and independent reviewers
Deployment	Excessive reliance on system recommendations	Meaningful trained human review	Welfare officers and supervisors
Decision Communication	Claimant cannot understand adverse outcome	Individualized, accessible explanation	Responsible public agency
Appeal	Automated errors become difficult to correct	Human reconsideration and evidence correction	Appeals authority
Monitoring	Performance deteriorates after deployment	Continuous monitoring and periodic independent audit	Agency and external oversight body
Retirement/Replacement	Harmful legacy data or rules continue unnoticed	Controlled decommissioning and record-retention policy	Agency governance board

The matrix shows that responsibility should remain institutionally identifiable even where technical functions are outsourced. A government agency may contract software development, infrastructure, or model maintenance, but it cannot logically outsource its public accountability to citizens. Vendor responsibility and agency responsibility should therefore be complementary rather than substitutive.

5.3 Illustrative Accountability Maturity Analysis

The following graph demonstrates how the seven accountability dimensions could be incorporated into an internal maturity assessment. The lower series represents an illustrative welfare system in which automation exists but formal accountability controls remain fragmented. The higher series shows a hypothetical target profile following implementation of the proposed framework.

Graph 1: Illustrative Accountability Maturity Profile for Technology-Enabled Welfare Administration



Note: Values are simulated for conceptual demonstration and must not be interpreted as empirical findings from an actual welfare authority.

The graph illustrates an important feature of the proposed model: accountability improvement should be multidimensional. Raising technical explainability while leaving appeals inaccessible would not create a fully accountable system. Similarly, strengthening data governance without independent auditability could leave institutional blind spots undetected. The maturity profile is therefore intended to reveal imbalances across safeguards rather than produce a single superficial compliance score.

6. Discussion

6.1 From Algorithm Transparency to Institutional Answerability

The analysis suggests that transparency should be understood as a means of accountability rather than its endpoint. Publishing source code, technical documentation, or model descriptions may assist specialized reviewers, but welfare claimants ordinarily require something different: understandable reasons for a decision, knowledge of how automation influenced that decision, and an effective procedure for correction.

Accountability therefore requires different levels of explanation for different audiences. Technical auditors may require model documentation, data lineage, error rates, and testing records. Administrative reviewers need information about relevant factors, confidence, exceptions, and permissible overrides. A claimant requires an accessible explanation of why the decision affected them and what evidence can be challenged. Regulators require information sufficient to evaluate legality, fairness, performance, and institutional compliance.

OECD's emphasis on traceability supports this layered conception by connecting accountability with the ability to examine datasets, processes, and decisions across the system lifecycle. Accountability



is consequently not equivalent to revealing every technical parameter. It is the ability to establish an evidence trail linking system design to administrative action and identifiable institutional responsibility.

6.2 Human Oversight Must Be Meaningful Rather Than Ceremonial

A major implication concerns the concept of the “human in the loop.” Human involvement should not be treated as proof of accountability unless the human reviewer possesses genuine decision authority. Research by Green challenges governance strategies that rely heavily on human oversight without examining whether humans can actually perform the expected corrective function. Sztandar-Sztanderska's welfare-administration research similarly indicates that the organizational context of frontline work significantly influences the practical capacity to supervise automated systems.

Meaningful review therefore requires institutional design. Caseworkers need training explaining what the algorithm predicts, what it does not predict, what data it uses, known limitations, and circumstances requiring escalation. Workload targets should not penalize justified departures from automated recommendations. Reviewers should be able to examine relevant supporting information rather than receiving only a risk category or recommendation.

Most importantly, agencies should identify decisions for which automation is inappropriate regardless of potential efficiency. Some welfare decisions involve contextual judgment, unusual hardship, conflicting evidence, or rights-sensitive circumstances that cannot be responsibly reduced to standardized predictive outputs. Algorithmic accountability therefore includes the capacity to decide not to automate.

6.3 Administrative Efficiency and Public Legitimacy

Technology-enabled welfare administration is frequently justified through efficiency, fraud reduction, consistency, and better targeting. These goals are legitimate, but administrative efficiency cannot be evaluated independently of error distribution. A system may reduce processing time while transferring the cost of mistakes to individuals least capable of enduring delayed payments, lengthy investigations, or complex appeals.

The emerging regulatory approach to high-risk AI reflects the principle that systems capable of materially affecting individuals require stronger safeguards. The EU AI Act combines accuracy and robustness obligations with oversight, registration, information, and monitoring requirements. The underlying governance lesson extends beyond Europe: high administrative efficiency does not by itself establish legitimacy.

For welfare agencies, trust is likely to depend on visible institutional willingness to identify and correct errors. An accountable organization does not claim that algorithms are infallible. It documents their limitations, monitors failure patterns, permits challenge, communicates reasons, and assigns responsibility for remediation. Public confidence is therefore more plausibly sustained through demonstrable contestability and institutional learning than through promises of technological objectivity.

7. Conclusion

Algorithmic systems are becoming increasingly important components of contemporary public administration, and their use in welfare systems presents both significant administrative opportunities



and substantial governance challenges. Automation can support faster processing, identify patterns in complex administrative data, assist resource allocation, and reduce certain repetitive burdens. Yet welfare administration is not an ordinary data-processing environment. Decisions frequently concern access to resources necessary for basic economic and social security, making accountability a fundamental institutional requirement.

This study has argued that algorithmic accountability should be understood as a lifecycle capacity rather than a narrow transparency requirement. Seven dimensions are central to this capacity: traceability, explainability, bias testing, meaningful human review, appeal and redress, data governance, and independent auditability. These safeguards must operate from system authorization and procurement through development, operational use, monitoring, and eventual retirement.

The analysis also demonstrates why the presence of a human official does not automatically make algorithmic administration accountable. Meaningful oversight requires authority, information, competence, time, and institutional support. Responsibility must remain identifiable at the organizational level even where private vendors or specialized technical teams participate in system development.

The principal contribution of the proposed framework is therefore to shift attention from a technology-centered question—Is the algorithm transparent?—toward an institution-centered question—Can the welfare administration explain, justify, review, correct, and accept responsibility for decisions produced with algorithmic assistance? This distinction is crucial. Public accountability is fulfilled not when an algorithm merely performs well, but when the public institution using it remains answerable for how technological power affects the rights, interests, and material security of citizens.

Future empirical research should test the framework in operational welfare agencies using administrative audits, beneficiary interviews, caseworker studies, appeal records, subgroup error analysis, and longitudinal performance monitoring. Comparative studies across jurisdictions could further determine which accountability mechanisms are most effective under different legal, organizational, and technological conditions.

References

1. Eubanks, V. (2018). *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor*. St. Martin's Press.
2. Brown, A., Chouldechova, A., Putnam-Hornstein, E., Tobin, A., & Vaithianathan, R. (2019). Toward algorithmic accountability in public services: A qualitative study of affected community perspectives on algorithmic decision-making in child welfare services. *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, 1–12.
3. Kleinberg, J., Ludwig, J., Mullainathan, S., & Sunstein, C. R. (2019). Discrimination in the age of algorithms. *Journal of Legal Analysis*, 10, 1–47.
4. Wirtz, B. W., Weyerer, J. C., & Geyer, C. (2019). Artificial intelligence and the public sector—Applications and challenges. *International Journal of Public Administration*, 42(7), 596–615.
5. Kroll, J. A., Huey, J., Barocas, S., Felten, E. W., Reidenberg, J. R., Robinson, D. G., & Yu, H. (2017). Accountable algorithms. *University of Pennsylvania Law Review*, 165(3), 633–705.
6. Pasquale, F. (2015). *The Black Box Society: The Secret Algorithms That Control Money and Information*. Harvard University Press.
7. Citron, D. K., & Pasquale, F. (2011). The scored society: Due process for automated predictions. *Washington Law Review*, 89, 1–33.



8. 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.
9. Selbst, A. D., & Barocas, S. (2018). The intuitive appeal of explainable machines. *Fordham Law Review*, 87(3), 1085–1139.
10. Burrell, J. (2016). How the machine “thinks”: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12.
11. Diakopoulos, N. (2015). Algorithmic accountability: Journalistic investigation of computational power structures. *Digital Journalism*, 3(3), 398–415.
12. Sandvig, C., Hamilton, K., Karahalios, K., & Langbort, C. (2014). Auditing algorithms: Research methods for detecting discrimination on internet platforms. *Data and Discrimination: Converting Critical Concerns into Productive Inquiry*, 1–23.
13. Veale, M., & Brass, I. (2019). Administration by algorithm? Public management meets public sector machine learning. In *Algorithmic Regulation*. Oxford University Press.
14. Janssen, M., & Kuk, G. (2016). The challenges and limits of big data algorithms in technocratic governance. *Government Information Quarterly*, 33(3), 371–377.
15. Wirtz, B. W., Weyerer, J. C., & Sturm, B. J. (2020). The dark sides of artificial intelligence: An integrated AI governance framework for public administration. *International Journal of Public Administration*, 43(9), 818–829.
16. Alston, P. (2019). *Report of the Special Rapporteur on Extreme Poverty and Human Rights: Digital Welfare States and Human Rights*. United Nations Human Rights Council. The report specifically addresses how digital technologies can be used to target, surveil, and administer welfare recipients and raises concerns about accountability and human rights.
17. Henman, P. (2019). *Governing Electronically: E-Government and the Reconfiguration of Public Administration*. Routledge.
18. Bovens, M., & Zouridis, S. (2002). From street-level to system-level bureaucracies: How information and communication technology is transforming administrative discretion and constitutional control. *Public Administration Review*, 62(2), 174–184.
19. Yeung, K. (2018). Algorithmic regulation: A critical interrogation. *Regulation & Governance*, 12(4), 505–523.
20. 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.